

PromptSculptor: Multi-Agent Based Text-to-Image Prompt Optimization

Dawei Xiang¹, Wenyan Xu², Kexin Chu¹, Tianqi Ding³,
Zixu Shen¹, Yiming Zeng¹, Jianchang Su¹, Wei Zhang^{1*},

¹University of Connecticut, ² Central University of Finance and Economics, ³ Baylor University

¹ {ieb24002, kexin.chu, qzt24001, yiming.zeng, jianchang.su, wei.13.zhang} @uconn.edu,

² 2022211032@email.cufe.edu.cn, ³ kirk_ding1@baylor.edu

Abstract

The rapid advancement of generative AI has democratized access to powerful tools such as Text-to-Image (T2I) models. However, to generate high-quality images, users must still craft detailed prompts specifying scene, style, and context—often through multiple rounds of refinement. We propose PromptSculptor, a novel multi-agent framework that automates this iterative prompt optimization process. Our system decomposes the task into four specialized agents that work collaboratively to transform a short, vague user prompt into a comprehensive, refined prompt. By leveraging Chain-of-Thought (CoT) reasoning, our framework effectively infers hidden context and enriches scene and background details. To iteratively refine the prompt, a self-evaluation agent aligns the modified prompt with the original input, while a feedback-tuning agent incorporates user feedback for further refinement. Experimental results demonstrate that PromptSculptor significantly enhances output quality and reduces the number of iterations needed for user satisfaction. Moreover, its model-agnostic design allows seamless integration with various T2I models, paving the way for industrial applications.

1 Introduction

The rapid development of large-scale models since 2022 has introduced Generative AI tools to a wide audience. Text-to-Image (T2I) models (Zhang et al., 2023) like Midjourney and Large Language Models (LLMs) such as ChatGPT now produce astonishing, often human-level creative outputs. However, effectively harnessing these technologies typically requires subtle and specific prompting. For T2I models, for instance, prompts must include precise details on scene composition, style parameters, and even technical jargon, creating a high entry barrier for new AI users (Mahdavi Goloujeh

et al., 2024). This challenge has led to the emergence of specialized professionals—"Midjourney Artists"—who need to optimize customers' brief inputs into detailed, high-quality prompts through iterative refinement based on generated outputs and user feedback.

Previous works (Cao et al., 2023; Feng et al., 2023; Mo et al., 2024) have tried to build a system to automate this prompt engineering workflow. But there are significant challenges for the system. A primary challenge is inferring user intent: initial prompts are often short and vague, peppered with abstract metaphors. For example, "draw a painting as a birthday blessing for my friend, he is like a lion" might lead a T2I model to interpret "lion" literally as a real fierce lion rather than as a symbol of confidence and courage. The second challenge is enriching these sparse inputs with detailed scene and background descriptions (Mahdavi Goloujeh et al., 2024). This requires the system to associate abstract concepts with concrete visual elements. Besides T2I models usually cannot fully satisfy user requirements in a single generation attempt. Therefore, prompts must be adjusted iteratively based on user feedback after observing initial results.

However previous works have primarily focused on parameter editing while neglecting to enrich prompts with detailed background and scene descriptions. As a result, the generated outputs often remain too similar to the original, lacking the necessary contextual depth. Moreover, these approaches typically do not address the challenge of interpreting abstract concepts—if a prompt contains hard-to-visualize terms, the generated image may fail to capture the user's true intent. In addition, most methods lack an iterative mechanism for updating the prompt based on feedback after the generation process.

To address these challenges, we propose a multi-agent system (MAS) that decomposes the task

*Corresponding author

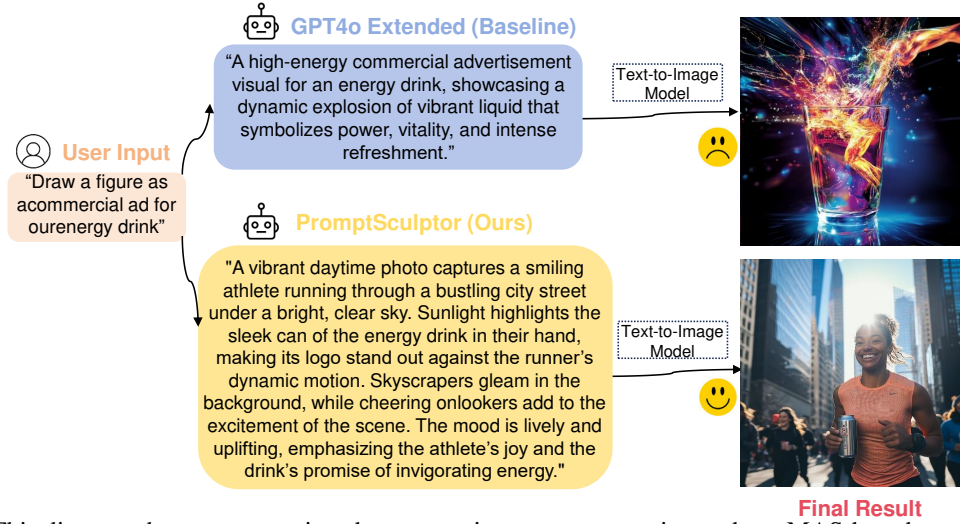


Figure 1: This diagram shows a comparison between naive prompt extension and our MAS-based prompt optimization. Our model successfully depicts a commercial ad scene to propaganda the drink which better aligns with user’s intent.

of prompt optimization into specialized functions. Our system includes four primary agents: 1) Intent Inference Agent; 2) Scene and Style Agent; 3) Feedback and Tuning Agent ; 4) Self-Evaluation Agent.

The core advantage of our method lies in its multi-agent architecture, which significantly enhances language understanding over traditional single-agent systems. First, our Intent Inference Agent decomposes brief and ambiguous inputs by breaking down abstract terms into detailed concepts and specific objects. This refined information is then passed to the Scene and Style Agent, which enriches the prompt with detailed scene, background, and style descriptions by linking these concepts to concrete visual elements. The Self-Evaluation Agent and the Feedback and Tuning Agent further refine the prompt until the final output meets expectations. Since our system is built on LLM agents and leverages user feedback, it can seamlessly transfer across different T2I models without requiring additional fine-tuning.

Our system also integrate Chain-of-Thought (CoT) reasoning (Wei et al., 2022) into user intent understanding and scene enrichment process, which requires agent to provides transparent, step-by-step rationales for each enrichment stage. This decomposition process increase the agent’s ability to understand complex terms and requirements.

Our innovation and contributions are as follows:

- We propose a novel multi-agent framework that decomposes the complex T2I prompt optimization task into specialized agents, each

handling a distinct subtask. This simplifies the process and significantly improves generation quality and flexibility compared to prior single-agent systems (Cao et al., 2023; Mo et al., 2024). To the best of our knowledge, this is the first use of a MAS for T2I prompt optimization, achieving state-of-the-art performance as shown in Table 1.

- Our system features a self-evaluation agent based on a Vision-Language Model (VLM) for semantic alignment, paired with a feedback-tuning agent. This feedback loop iteratively refines prompts until the generated image matches the user’s intent, reducing the number of required runs, as demonstrated in Table 2. This is the first introduction of such a feedback mechanism in T2I prompt optimization.
- Our framework is model-agnostic and works across various T2I models (e.g., Midjourney, DALL-E 3, Stable Diffusion) without model-specific tuning, ensuring strong scalability and broad applicability.

2 Related Work

2.1 Prompt optimization

Prompt optimization has gained attention for enhancing LLMs without parameter fine-tuning. Early research focused on white-box models like AutoPrompt(Shin et al., 2020) and Fluent-Prompt(Shi et al.), which optimize by accessing embeddings or logits. With the rise of closed-

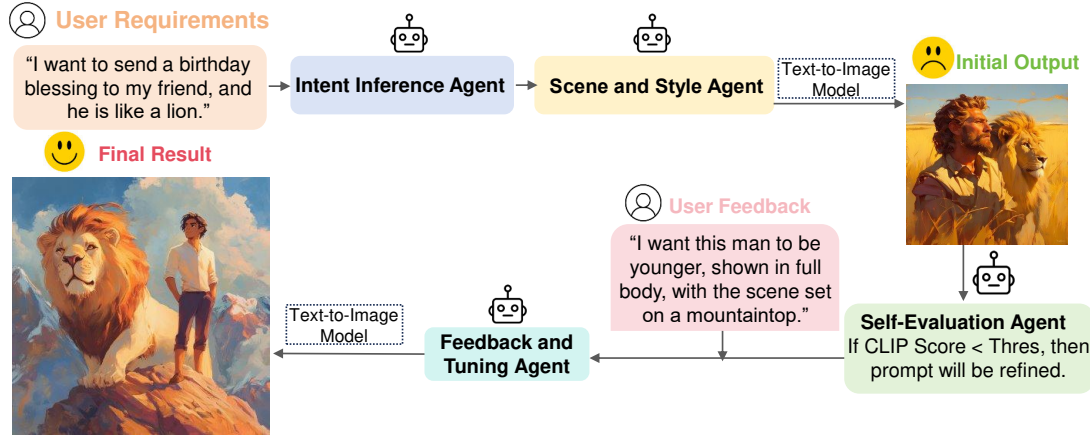


Figure 2: This diagram shows the PromptSculptor model’s process, where the user’s prompt is analyzed, enriched, and refined by different agents. The system optimizes the prompt and generates an image that aligns with the user’s intent, adjusting based on feedback and evaluation scores.

source models like GPT(Brown et al., 2020), research shifted toward black-box optimization. While methods like BBTv2(Sun et al., 2022), and Clip-Tuning still rely on embeddings, recent approaches like GRIPS(Prasad et al., 2023) and APO(Pryzant et al., 2023) optimize by editing and selecting candidate prompts. Other methods use evolutionary algorithms, reinforcement learning, or planning strategies, such as EvoPrompt(Guo et al., 2025), Promptbreeder(Fernando et al., 2024), BDPL(Diao et al.), and PromptAgent(Wang et al.). APE(Zhou et al.) generates candidate prompts with LLMs, and OPRO(Zhang et al., 2024) frames optimization as a black-box problem solved by LLMs. OPT2I(Mañas et al.) refines prompts to improve consistency. Some studies, like InstructZero(Chen et al., 2024b) and INSTINCT(Lin et al.), treat prompt optimization as a continuous problem, while ZOPO(Hu et al., 2024) enhances performance through zero-order optimization. Recently, DPO(Rafailov et al., 2023) and SLiC(Khorasgani et al., 2022) bypass reinforcement learning using preference datasets. Unlike these methods, this paper proposes a plug-and-play multi-agent system for prompt optimization, without fine-tuning.

2.2 Multi-agent system

Multi-agent system works by coordinating multiple large language models to finish a complex task. Recent advances in large language models have catalyzed multi-agent paradigms into dynamic coordination (Chen et al., 2024a; Guo et al., 2024; Leong and Wu, 2024), which has being explored

in multiple areas including text and code generation (Zhang et al., 2025b; Zeng et al., 2025b), data labeling (Lu et al., 2025), bioinformatics analysis (Wu et al., 2024), financial forecasting(Xu et al., 2025b,c,a; Zhang et al., 2025a), autonomous driving (Zeng et al., 2025a) and privacy protection (Chu et al., 2025b,a). Building on these trends, we position our approach as a plug-and-play multi-agent prompt optimization framework that operates at the prompt level without fine-tuning, and is compatible with closed-source or API-only models.

2.3 Improve the Consistency of T2I Model

In recent years, some studies have focused on using cross-attention to guide LLMs in generating prompts that better align with user intent(Feng et al.; Epstein et al., 2023; Wu et al., 2023a; Chefer et al., 2023). Other research transforms text prompts into layouts, which are then fed into layout-to-image generation models(Cho et al., 2023; Lian et al.). Some approaches fine-tune T2I models using human feedback(Lee et al., 2023; Wu et al., 2023b; Wallace et al., 2024), AI model feedback(Sun et al.), or image selection(Karthik et al., 2023). Unlike these methods, our approach doesn’t require model fine-tuning but uses a correction agent to check if the input prompt aligns with the natural language query, operating only at the prompt level. This allows our method to work with more T2I models, including those only accessible via API as Midjourney (Midjourney, 2024).

3 Methodology

In this section, we introduce the design of our model, PromptSculptor, which comprises three main stages: (1) a MAS for initial prompt optimization, (2) VLM-based prompt alignment and (3) a feedback-based tuning agent. Our system is intended for users with limited experience in prompt engineering, where the initial T2I prompt may be vague or too simplistic. To address this, our system first analyzes the user's intent to automatically enrich and detail the prompt. It then evaluates whether the generated image aligns with the user's intent and, if necessary, further refines the prompt using a pre-trained VLM.

3.1 Multi-Agent System Overview

As shown in Figure 2, our system consists of four specialized agents:

1. **Intent Inference Agent:** Extracts the user's core idea and missing details.
2. **Scene and Style Agent:** Builds on that intent to craft a vivid, detailed scene.
3. **Self-Evaluation Agent:** Generates an image, checks its CLIP score against the original intent, and fine-tunes the prompt if the score is too low.
4. **Feedback and Tuning Agent:** Incorporates user feedback to further refine the prompt and close any remaining gaps.

3.2 Intent Inference Agent

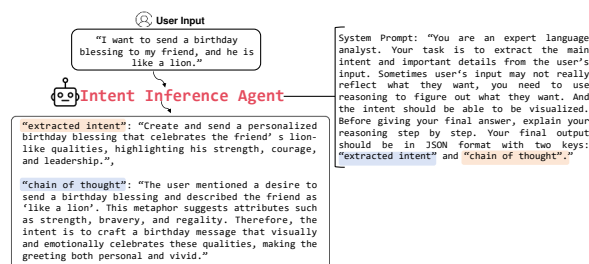


Figure 3: Prompt display of the Intent Inference Agent.

The primary function of the Intent Inference Agent is to perform a deep analysis of the user's input, capturing not only the explicit request but also the latent, nuanced intent embedded within brief or ambiguous prompts. As shown in Figure 3, our system prompt is meticulously designed to guide the model to probe beyond surface-level

text—extracting implicit cues, contextual hints, and even emotional undertones.

To ensure transparency and enhance interpretability, the agent is equipped with a CoT mechanism. This mechanism prompts the agent to document its reasoning process step by step. For example, when processing a prompt like "A birthday painting for a friend who is like a lion," the agent's CoT may detail steps such as identifying the explicit elements ("birthday painting" and "friend"), interpreting "lion" as a metaphor for qualities like strength, majesty, or courage, and synthesizing these insights to reveal the underlying intent. This comprehensive extraction of both overt and hidden intentions, supported by its explicit CoT, forms the foundation for further prompt enrichment in our multi-agent pipeline.

3.3 Scene and Style Agent

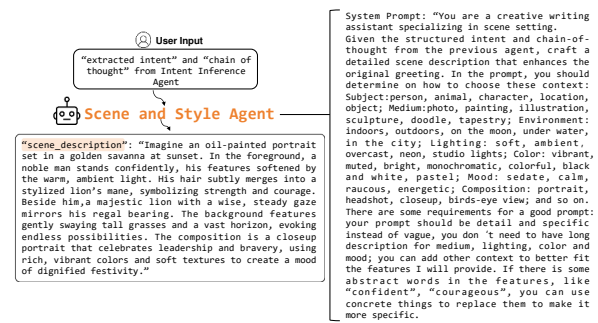


Figure 4: Prompt display of the Scene and Style Agent.

After Intent Inference Agent has detailed the user's intent and context, the second step is for Scene and Style Agent to set up the scene and enrich the detail in the figure. As shown in Figure 4 this agent will try to visualize the abstract concept in the previous prompt and use detail object to visualize them to make the figure more vivid. This is exactly what human artists are doing: they use some detailed object to represent abstract things to make them more impressive and easy to understand.

When artist creates a painting, there are several factors they need to consider including the subjects, medium, environment, lighting, color, mood composition and so on. We all instruct this agent to consider all these factors based on the previous inference. If without the inference and context explanation from the previous agent, it may be more difficult for the Scene and Style Agent to generate an accurate settings. This shows the advantage of our MAS.

3.4 Self-Evaluation agent

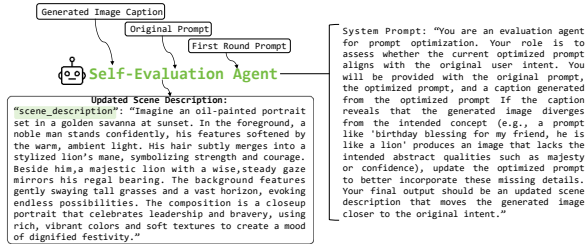


Figure 5: Prompt display of the Self-Evaluation Agent.

Although the Scene and Style Agent enriches the scene description, the inherent vagueness of the original input may still result in generated images that do not fully align with the user’s intent. To address this gap, we designed a self-evaluation agent to further refine the prompt if needed.

Our self-evaluation agent serves as a quality assurance module within our multi-agent prompt optimization pipeline. Its primary role is to verify that the generated image accurately reflects the user’s original intent. To achieve this, the agent first computes a CLIP similarity score (Patashnik et al., 2021) between the generated image and the original prompt. As shown in Algorithm 1 if the similarity score falls below a predetermined threshold, the agent leverages the BLIP-2 model (Li et al., 2023)—a state-of-the-art VLM—to generate a detailed caption for the image. This caption provides an independent description of the image’s visual content.

By comparing the BLIP-2-generated caption with both the original prompt and the current optimized prompt, the agent identifies discrepancies or missing elements. For example, if the user’s input emphasizes abstract qualities like “lion-like majesty” but the caption describes only a generic birthday scene, the agent infers that critical visual attributes (e.g., strength, regality) may be under-represented. Based on this analysis, the system automatically refines the optimized prompt, adding or modifying details to better capture the intended semantic content.

This iterative feedback loop—driven by the combination of CLIP-based evaluation and BLIP-2 captioning—ensures that the enriched prompt aligns closely with the user’s intent. It provides a robust, quantitative measure to guide the continuous improvement of prompt quality, ultimately enabling consistent and high-quality image generation across different T2I models.

Algorithm 1 Self-Evaluation Agent (SEA) for Prompt Improvement

Require: Generated Image I , Original Prompt P_o ,

Optimized Prompt P_{opt} , Threshold τ

- 1: $s \leftarrow \text{CLIP}(I, P_o)$ ▷ Compute similarity score
- 2: **if** $s \geq \tau$ **then**
- 3: **return** P_{opt} ▷ No further refinement needed
- 4: **else**
- 5: $C \leftarrow \text{BLIP2}(I)$ ▷ Generate a caption for the image
- 6: $P_{improved} \leftarrow \text{SEA}(P_o, P_{opt}, C)$ ▷ Refine prompt
- 7: **return** $P_{improved}$
- 8: **end if**

3.5 Feedback and Tuning Agent

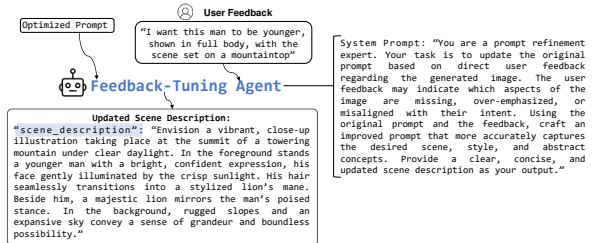


Figure 6: Prompt display of the Feedback-Tuning Agent.

While the Self-Evaluation Agent ensures that the generated image scores highly against the original prompt, it may still miss the user’s true intent. To bridge this gap, we introduce a Feedback and Tuning Agent that takes direct user feedback and uses it to iteratively refine the prompt. Because the initial prompt often contains only limited information, this feedback loop is essential for aligning the final image with the user’s vision. As shown in Figure 6, the agent applies user comments to continuously improve the optimized prompt until the generated result fully meets the user’s requirements.

4 Experiments

4.1 Experiment Settings

Implementation detail We used the API from OpenAI to build our multi-agent system. We use GPT-4o as the basemodel. Each input prompt will first pass the Intent Inference Agent and Scene and Style Agent sequentially. Then an image will generate from the extended prompt. Later the Self-Evaluation agent will keep improve the prompt until its generated image is aligned with the original prompt. If there are feedbacks from user, the Feedback and Tuning Agent will further improve the prompt. We utilize the Stable-Diffusion-

Method	CLIP Score(↑)	PickScore(↑)	Aes Score(↑)
Original	0.289	19.43	5.87
Extended	0.232	20.28	6.21
MagicPrompt	0.246	18.69	6.11
Ours w/o SEA	0.257	20.26	6.68
Ours	0.263	21.31	6.96

Table 1: Performance comparison for different methods. "SEA" stands for Self-Evaluation Agent.

XL(SDXL) (Podell et al., 2023) and Midjourney (Midjourney, 2024) as the T2I model. For each test, all methods to compare will choose the same T2I model.

Methods to compare We have selected a few famous models as our baselines and counterparts:

- Original prompt.
- Simple extend prompt: we ask GPT 4 to add more detail to the original prompt.
- MagicPrompt (Gustavosta): A prompt enrichment framework trained with 80000 images from stable diffusion. They use simple-enriched prompt pair to fine-tune a GPT-2.
- PromptAgent(Wang et al.): State-of-the-Art in prompt optimization using multi-agent system.

Metrics It is not straightforward to compare the quality of different optimized prompts. We use the generated image from T2I model as a metrics for the quality of a prompt. Then we calculate the PickScore (Kirstain et al., 2023), CLIP score (Radford et al., 2021), Aesthetic score (Schuhmann et al., 2022) between original prompts and the generated images. CLIP score evaluate the similarity between the prompts and the generated image. PickScore evaluate the alignment between the prompts and the generated image and also the overall quality. Aesthetic score evaluate the human preference on the aesthetic of the generated images.

4.2 Overall Results

Table 1 presents a comprehensive performance comparison. Here, "Original" denotes the initial short prompt, "Extended" refers to the GPT-4 expanded prompt, and "MagicPrompt" or "PromptAgent" represents the prompt extended using MagicPrompt or PromptAgent from the original simple prompt. Our method achieves the highest PickScore, Aesthetic Score and human expert preference score, demonstrating that it effectively aligns the optimized prompt with the generated

Method	Preference Score(↑)	Number of Runs(↓)
Original	69.85 %	6.08
Extended	75.32 %	4.22
MagicPrompt	67.28 %	5.33
Ours	80.12 %	2.35

Table 2: Human Evaluation Comparison.

image and produces higher-quality images with enhanced aesthetic appeal and user preference. Additionally, our method attains the second highest CLIP score, indicating strong alignment between the generated image and the original user request. Note that the CLIP score is calculated using the original prompt and the generated image, which is why the image generated from the original prompt shows the highest similarity to that prompt.

Ablation Study Table 1 shows that our Self-Evaluation Agent (SEA) consistently improve the evaluation performance compared to without this agent. Thanks to the automatically alignment evaluation, our method can achieve higher output quality without additional user’s feedback.

4.3 Human Evaluation

In addition to our model-based evaluation, we conducted a human evaluation. Twenty volunteers generated 60 real-world prompts that were intentionally vague, using abstract words and metaphors. They then rated the first-round generated images on a scale from 0 to 100 and recorded the number of prompt modifications needed to reach satisfaction. As shown in Table 2, our method consistently achieved the highest preference scores. This indicates that our approach meets real-world prompt auto-completion requirements. Moreover, thanks to the Self-Evaluation Agent’s ability to identify misalignments and the Feedback-Tuning Agent’s efficient prompt optimization, our method required the fewest modifications to achieve satisfaction.

4.4 Visualization Results

We have included visualization results in the Appendix A. Our method consistently shows better output quality and close alignment to users’ intent compared to other methods.

We also include the prompt comparison in the appendix.

5 Industry Impact

We are collaborating with a startup to build a platform that integrates T2I model prompt auto-completion and optimization. As we have mentioned, there are many people who lack Prompt Engineering Experience but are willing to try Text-to-Image models. The goal of this platform is to empower these users in generating impressive figures from just a simple idea. The platform is under development and is supposed to be public very soon.

6 Conclusion

In this paper, we introduce PromptSculptor, a novel multi-agent framework that automates the iterative refinement of T2I prompts. By decomposing the complex prompt refinement process into specialized agents, our approach produces detailed, context-aware prompts that better align with user expectations from simple and vague input. Our experiments demonstrate that PromptSculptor significantly enhances both the quality and aesthetic appeal of generated images from simple inputs.

References

- Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, et al. 2020. Language models are few-shot learners. *Advances in neural information processing systems*, 33:1877–1901.
- Tingfeng Cao, Chengyu Wang, Bingyan Liu, Ziheng Wu, Jinhui Zhu, and Jun Huang. 2023. [Beautiful Prompt: Towards automatic prompt engineering for text-to-image synthesis](#). In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing: Industry Track*, pages 1–11, Singapore. Association for Computational Linguistics.
- Hila Chefer, Yuval Alaluf, Yael Vinker, Lior Wolf, and Daniel Cohen-Or. 2023. Attend-and-excite: Attention-based semantic guidance for text-to-image diffusion models. *ACM transactions on Graphics (TOG)*, 42(4):1–10.
- Junzhe Chen, Xuming Hu, Shuodi Liu, Shiyu Huang, Wei-Wei Tu, Zhaofeng He, and Lijie Wen. 2024a. Llmarena: Assessing capabilities of large language models in dynamic multi-agent environments. *arXiv preprint arXiv:2402.16499*.
- Lichang Chen, Jiuhai Chen, Tom Goldstein, Heng Huang, and Tianyi Zhou. 2024b. Instructzero: Efficient instruction optimization for black-box large language models. In *International Conference on Machine Learning*, pages 6503–6518. PMLR.
- Jaemin Cho, Abhay Zala, and Mohit Bansal. 2023. Visual programming for step-by-step text-to-image generation and evaluation. *Advances in Neural Information Processing Systems*, 36:6048–6069.
- Kexin Chu, Zecheng Lin, Dawei Xiang, Zixu Shen, Jianchang Su, Cheng Chu, Yiwei Yang, Wenhui Zhang, Wenfei Wu, and Wei Zhang. 2025a. Selective kv-cache sharing to mitigate timing side-channels in llm inference. *arXiv preprint arXiv:2508.08438*.
- Kexin Chu, Zixu Shen, Dawei Xiang, and Wei Zhang. 2025b. Safekv: Safe kv-cache sharing in llm serving. In *Machine Learning for Computer Architecture and Systems*.
- Shizhe Diao, Zhichao Huang, Ruijia Xu, Xuechun Li, LIN Yong, Xiao Zhou, and Tong Zhang. Black-box prompt learning for pre-trained language models. *Transactions on Machine Learning Research*.
- Dave Epstein, Allan Jabri, Ben Poole, Alexei Efros, and Aleksander Holynski. 2023. Diffusion self-guidance for controllable image generation. *Advances in Neural Information Processing Systems*, 36:16222–16239.
- Weixi Feng, Xuehai He, Tsu-Jui Fu, Varun Jampani, Arjun Reddy Akula, Pradyumna Narayana, Sugato Basu, Xin Eric Wang, and William Yang Wang. Training-free structured diffusion guidance for compositional text-to-image synthesis. In *The Eleventh International Conference on Learning Representations*.
- Yingchaojie Feng, Xingbo Wang, Kam Kwai Wong, Si-jia Wang, Yuhong Lu, Minfeng Zhu, Baicheng Wang, and Wei Chen. 2023. Promptmagician: Interactive prompt engineering for text-to-image creation. *IEEE Transactions on Visualization and Computer Graphics*, 30(1):295–305.
- Chrisantha Fernando, Dylan Sunil Banarse, Henryk Michalewski, Simon Osindero, and Tim Rocktäschel. 2024. Promptbreeder: Self-referential self-improvement via prompt evolution. In *International Conference on Machine Learning*, pages 13481–13544. PMLR.
- Taicheng Guo, Xiuying Chen, Yaqi Wang, Ruidi Chang, Shichao Pei, Nitesh V Chawla, Olaf Wiest, and Xi-angliang Zhang. 2024. Large language model based multi-agents: A survey of progress and challenges. *arXiv preprint arXiv:2402.01680*.
- Ziyu Guo, Renrui Zhang, Chengzhuo Tong, Zhizheng Zhao, Peng Gao, Hongsheng Li, and Pheng-Ann Heng. 2025. Can we generate images with cot? let’s verify and reinforce image generation step by step. *arXiv preprint arXiv:2501.13926*.

- Gustavosta. [Gustavosta/magicprompt-stable-diffusion](#).
- Wenyang Hu, Yao Shu, Zongmin Yu, Zhaoxuan Wu, Xiaoqiang Lin, Zhongxiang Dai, See-Kiong Ng, and Bryan Kian Hsiang Low. 2024. Localized zeroth-order prompt optimization. *Advances in Neural Information Processing Systems*, 37:86309–86345.
- Shyamgopal Karthik, Karsten Roth, Massimiliano Mancini, and Zeynep Akata. 2023. If at first you don't succeed, try, try again: Faithful diffusion-based text-to-image generation by selection. *arXiv preprint arXiv:2305.13308*.
- Salar Hosseini Khorasgani, Yuxuan Chen, and Florian Shkurti. 2022. Slic: Self-supervised learning with iterative clustering for human action videos. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 16091–16101.
- Yuval Kirstain, Adam Polyak, Uriel Singer, Shahbuland Matiana, Joe Penna, and Omer Levy. 2023. Pick-a-pic: An open dataset of user preferences for text-to-image generation. *Advances in Neural Information Processing Systems*, 36:36652–36663.
- Kimin Lee, Hao Liu, Moonkyung Ryu, Olivia Watkins, Yuqing Du, Craig Boutilier, Pieter Abbeel, Mohammad Ghavamzadeh, and Shixiang Shane Gu. 2023. Aligning text-to-image models using human feedback. *arXiv preprint arXiv:2302.12192*.
- H.Y. Leong and Y. Wu. 2024. [Why should next-gen llm multi-agent systems move beyond fixed architectures to dynamic, input-driven graphs?](#) *SSRN Electronic Journal*.
- Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. 2023. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. In *International conference on machine learning*, pages 19730–19742. PMLR.
- Long Lian, Boyi Li, Adam Yala, and Trevor Darrell. Llm-grounded diffusion: Enhancing prompt understanding of text-to-image diffusion models with large language models. *Transactions on Machine Learning Research*.
- Xiaoqiang Lin, Zhaoxuan Wu, Zhongxiang Dai, Wenyang Hu, Yao Shu, See-Kiong Ng, Patrick Jaillet, and Bryan Kian Hsiang Low. Use your instinct: Instruction optimization for llms using neural bandits coupled with transformers. In *Forty-first International Conference on Machine Learning*.
- Yao Lu, Zhaiyuan Ji, Jiawei Du, Yu Shanqing, Qi Xuan, and Tianyi Zhou. 2025. From llm-ation to llm-orchestrator: Coordinating small models for data labeling. *arXiv preprint arXiv:2506.16393*.
- Atefeh Mahdavi Goloujeh, Anne Sullivan, and Brian Magerko. 2024. Is it ai or is it me? understanding users' prompt journey with text-to-image generative ai tools. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*, pages 1–13.
- Oscar Mañas, Pietro Astolfi, Melissa Hall, Candace Ross, Jack Urbanek, Adina Williams, Aishwarya Agrawal, Adriana Romero-Soriano, and Michal Drozdal. Improving text-to-image consistency via automatic prompt optimization. *Transactions on Machine Learning Research*.
- Midjourney. 2024. Midjourney. <https://www.midjourney.com/>. Accessed: 2024-07-05.
- Wenyi Mo, Tianyu Zhang, Yalong Bai, Bing Su, Ji-Rong Wen, and Qing Yang. 2024. Dynamic prompt optimizing for text-to-image generation. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 26627–26636.
- Or Patashnik, Zongze Wu, Eli Shechtman, Daniel Cohen-Or, and Dani Lischinski. 2021. Styleclip: Text-driven manipulation of stylegan imagery. In *Proceedings of the IEEE/CVF international conference on computer vision*, pages 2085–2094.
- Dustin Podell, Zion English, Kyle Lacey, Andreas Blattmann, Tim Dockhorn, Jonas Müller, Joe Penna, and Robin Rombach. 2023. Sdxl: Improving latent diffusion models for high-resolution image synthesis. *arXiv preprint arXiv:2307.01952*.
- Archiki Prasad, Peter Hase, Xiang Zhou, and Mohit Bansal. 2023. Grips: Gradient-free, edit-based instruction search for prompting large language models. In *EACL*.
- Reid Pryzant, Dan Iter, Jerry Li, Yin Tat Lee, Chengguang Zhu, and Michael Zeng. 2023. Automatic prompt optimization with "gradient descent" and beam search. *arXiv preprint arXiv:2305.03495*.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PmLR.
- Rafael Rafailov, Archit Sharma, Eric Mitchell, Christopher D Manning, Stefano Ermon, and Chelsea Finn. 2023. Direct preference optimization: Your language model is secretly a reward model. *Advances in Neural Information Processing Systems*, 36:53728–53741.
- Christoph Schuhmann, Romain Beaumont, Richard Vencu, Cade Gordon, Ross Wightman, Mehdi Cherti, Theo Coombes, Aarush Katta, Clayton Mullis, Mitchell Wortsman, et al. 2022. Laion-5b: An open large-scale dataset for training next generation image-text models. *Advances in neural information processing systems*, 35:25278–25294.
- Weijia Shi, Xiaochuang Han, Hila Gonen, Ari Holtzman, Yulia Tsvetkov, and Luke Zettlemoyer. Toward human readable prompt tuning: Kubrick's the shining is a good movie, and a good prompt too? In *The 2023 Conference on Empirical Methods in Natural Language Processing*.

- Taylor Shin, Yasaman Razeghi, Robert L Logan IV, Eric Wallace, and Sameer Singh. 2020. Autoprompt: Eliciting knowledge from language models with automatically generated prompts. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 4222–4235.
- Jiao Sun, Deqing Fu, Yushi Hu, Su Wang, Royi Rassin, Da-Cheng Juan, Dana Alon, Charles Herrmann, Sjoerd van Steenkiste, Ranjay Krishna, et al. Dreamsync: Aligning text-to-image generation with image understanding feedback. In *Synthetic Data for Computer Vision Workshop @ CVPR 2024*.
- Tianxiang Sun, Zhengfu He, Hong Qian, Yunhua Zhou, Xuan-Jing Huang, and Xipeng Qiu. 2022. Bbtv2: Towards a gradient-free future with large language models. In *Proceedings of the 2022 Conference on Empirical Methods in Natural Language Processing*, pages 3916–3930.
- Bram Wallace, Meihua Dang, Rafael Rafailov, Linqi Zhou, Aaron Lou, Senthil Purushwalkam, Stefano Ermon, Caiming Xiong, Shafiq Joty, and Nikhil Naik. 2024. Diffusion model alignment using direct preference optimization. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 8228–8238.
- Xinyuan Wang, Chenxi Li, Zhen Wang, Fan Bai, Hao-tian Luo, Jiayou Zhang, Nebojsa Jojic, Eric Xing, and Zhiting Hu. Promptagent: Strategic planning with language models enables expert-level prompt optimization. In *The Twelfth International Conference on Learning Representations*.
- Jason Wei, Xuezhi Wang, Dale Schuurmans, Maarten Bosma, Fei Xia, Ed Chi, Quoc V Le, Denny Zhou, et al. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35:24824–24837.
- Qiucheng Wu, Yujian Liu, Handong Zhao, Trung Bui, Zhe Lin, Yang Zhang, and Shiyu Chang. 2023a. Harnessing the spatial-temporal attention of diffusion models for high-fidelity text-to-image synthesis. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 7766–7776.
- Xiaoshi Wu, Keqiang Sun, Feng Zhu, Rui Zhao, and Hongsheng Li. 2023b. Human preference score: Better aligning text-to-image models with human preference. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 2096–2105.
- Xidong Wu, Sumin Jo, Yiming Zeng, Arun Das, Ting-He Zhang, Parth Patel, Yuanjing Wei, Lei Li, Shou-Jiang Gao, Jianqiu Zhang, et al. 2024. regulogpt: Harnessing gpt for end-to-end knowledge graph construction of molecular regulatory pathways. In *2024 IEEE EMBS International Conference on Biomedical and Health Informatics (BHI)*, pages 1–8. IEEE.
- Wenyan Xu, Rundong Wang, Chen Li, Yonghong Hu, and Zhonghua Lu. 2025a. Hrft: Mining high-frequency risk factor collections end-to-end via transformer. In *Companion Proceedings of the ACM on Web Conference 2025*, pages 538–547.
- Wenyan Xu, Dawei Xiang, Yue Liu, Xiyu Wang, Yanxiang Ma, Liang Zhang, Chang Xu, and Jiaheng Zhang. 2025b. Finmultitime: A four-modal bilingual dataset for financial time-series analysis. *arXiv preprint arXiv:2506.05019*.
- Wenyan Xu, Dawei Xiang, Rundong Wang, Yonghong Hu, Liang Zhang, Jiayu Chen, and Zhonghua Lu. 2025c. Learning explainable stock predictions with tweets using mixture of experts. *arXiv preprint arXiv:2507.20535*.
- Shuang Zeng, Xinyuan Chang, Mengwei Xie, Xinran Liu, Yifan Bai, Zheng Pan, Mu Xu, and Xing Wei. 2025a. Futuresightdrive: Thinking visually with spatio-temporal cot for autonomous driving. *arXiv preprint arXiv:2505.17685*.
- Yiming Zeng, Wanhao Yu, Zexin Li, Tao Ren, Yu Ma, Jinghan Cao, Xiyan Chen, and Tingting Yu. 2025b. Bridging the editing gap in llms: Fineedit for precise and targeted text modifications. *Preprint, arXiv:2502.13358*.
- Chenshuang Zhang, Chaoning Zhang, Mengchun Zhang, and In So Kweon. 2023. Text-to-image diffusion models in generative ai: A survey. *arXiv preprint arXiv:2303.07909*.
- J. Zhang, J. Gao, W. Ouyang, W. Zhu, and H.Y. Leong. 2025a. Time-llama: Adapting large language models for time series modeling via dynamic low-rank adaptation. In *Proceedings of the 63rd Annual Meeting of the Association for Computational Linguistics (Volume 4: Student Research Workshop)*. Association for Computational Linguistics (ACL 2025). Poster.
- Tuo Zhang, Jinyue Yuan, and Salman Avestimehr. 2024. Revisiting opro: The limitations of small-scale llms as optimizers. In *Findings of the Association for Computational Linguistics ACL 2024*, pages 1727–1735.
- Zhenhua Zhang, Jianfeng Wang, Zhengyang Li, Yunpeng Wang, and Jiayun Zheng. 2025b. Anncoder: A mti-agent-based code generation and optimization model. *Symmetry*, 17(7).
- Yongchao Zhou, Andrei Ioan Muresanu, Ziwen Han, Keiran Paster, Silviu Pitis, Harris Chan, and Jimmy Ba. Large language models are human-level prompt engineers. In *The Eleventh International Conference on Learning Representations*.

A Appendix

Figure 7 presents four prompt types used to generate images for six abstract themes. The results for each theme are analyzed as follows:

In Figure 8, our prompt shows a schoolboy imagining himself as a future astronaut exploring space,

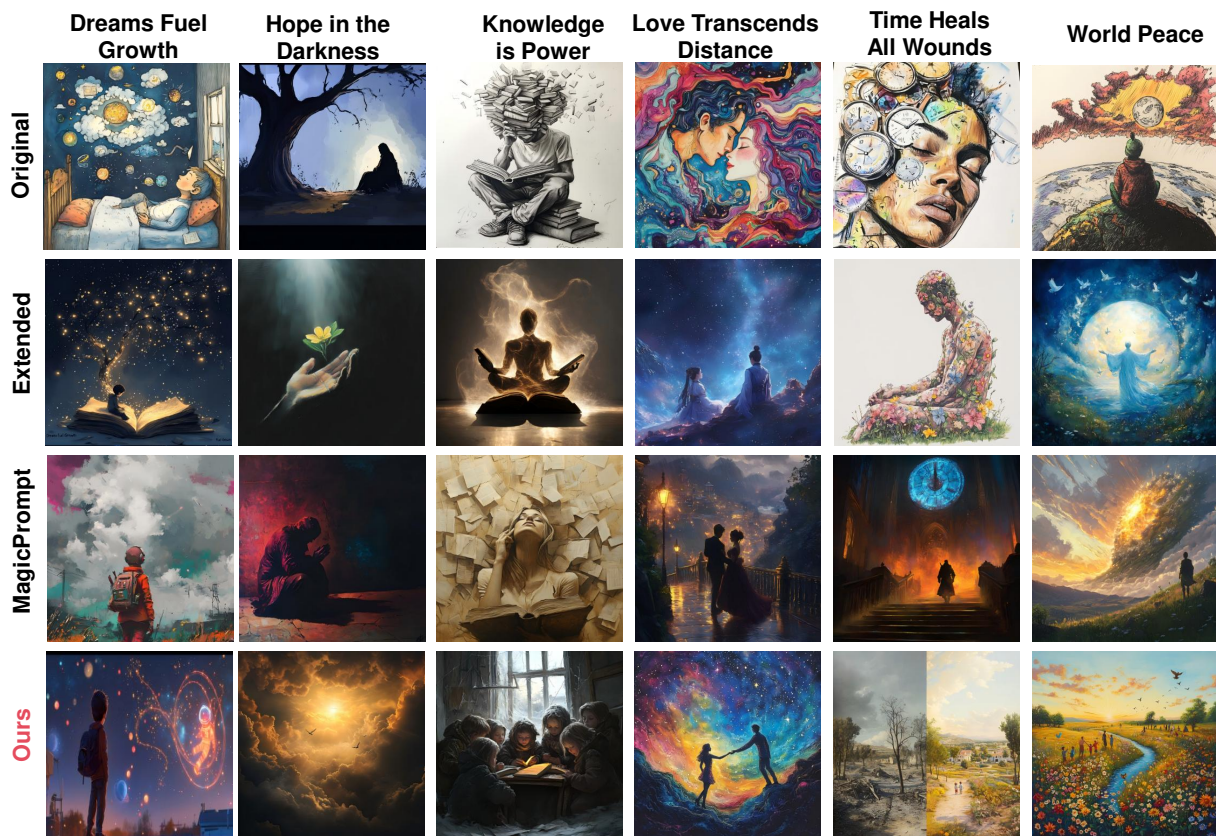


Figure 7: Four Prompt Types for Constructing Images Based on Six Abstract Themes

creating a vivid visual metaphor that highlights the theme “Dreams Fuel Growth.” In contrast, other prompts are more abstract or simplistic, lacking the kind of concrete and engaging imagery that sparks imagination. Moving to Figure 9, our prompt depicts sunlight breaking through dark clouds, with doves symbolizing hope, freedom, and resilience. This scene conveys a deeper emotional meaning and a sense of uplift, whereas the other prompts fall short of capturing such layered symbolism. In Figure 10, our prompt illustrates children in a rural classroom studying diligently, expressing their desire to change their fate through knowledge. This powerful imagery embodies hope and perseverance in the face of hardship. By comparison, the other prompts rely on more literal depictions of reading and fail to capture the abstract essence behind the pursuit of education. However, in Figure 11, none of the prompts—including ours—clearly emphasize the crucial sense of distance needed to effectively express the idea that “True Love Transcends Distance.” This remains an area for improvement. In Figure 12, unlike other prompts that directly show injured people and clocks, our prompt uses the imagery of rebuilding after a fire to symbolically convey the idea that “Time Heals

All Wounds.” This metaphorical approach offers a more thoughtful and layered interpretation. Finally, in Figure 13, our prompt integrates more human and natural elements than the others, with greater attention to emotional depth and visual detail. As a result, it presents a more genuine and optimistic atmosphere of peace, effectively capturing the intended theme.

In summary, our prompts offer a more effective interpretation and representation of abstract concepts.

Theme / By method	Dreams Fuel Growth
Original Prompts	I want to draw a figure showing “Dreams Fuel Growth”.
Extended Prompts	I want to draw a figure showing “Dreams Fuel Growth” – a child gazing up at a glowing tree sprouting from an open book, its branches filled with stars and symbols of imagination, reaching toward the sky.”
Magic Prompts	I want to draw a figure titled “Dreams Fuel Growth”, in the style of +250k concept art and character illustrations by Greg Rutkowski, Craig Mullins, and Daniel Dociu – trending on ArtStation. The scene features a cloudy background and ultra-realistic digital art.
Our Prompts	I want to draw a figure showing “Dreams Fuel Growth”. A young boy with a backpack stands on a quiet rooftop at dusk, gazing up at the sky. On the right side of the scene, his imagined future self floats in space as an astronaut, surrounded by stars and planets. A glowing trail of stardust connects the boy to his dream, symbolizing growth and aspiration. Soft ambient lighting, vibrant colors against a deep twilight sky, calm and inspiring mood, cinematic composition.”.

Figure 8: Four Prompt Types for Constructing Images Based on the Theme 'Dreams Fuel Growth'

Theme / By method	Hope in the Darkness
Original Prompts	I want to draw a figure showing “Hope in the Darkness”.
Extended Prompts	I want to draw a figure showing "Hope in the Darkness" – a single light source breaking through shadows, highlighting a small symbol of resilience like a blooming flower or an outstretched hand.
Magic Prompts	“I want to draw a color figure showing "Hope in the Darkness" (W1 9 8 2), trending on artstation, award winning painting, cgi, art by greg rutkowski
Our Prompts	I want to draw a figure showing “Hope in the Darkness”. A breathtaking scene above dense, stormy black clouds – a radiant golden sun breaks through, casting light that the darkness cannot reach. High above, a few birds glide effortlessly in the sunlight, symbolizing freedom and resilience. The contrast between the shadowy clouds below and the glowing sky above creates a powerful image of hope rising beyond despair.

Figure 9: Four Prompt Types for Constructing Images Based on the Theme 'Hope in the Darkness'

Theme / By method	Knowledge is Power
Original Prompts	I want to draw a figure showing “Knowledge is Power”.
Extended Prompts	I want to draw a figure showing “Knowledge is Power” – a stylized silhouette absorbing light from an open book, radiating energy or transformation, in a clean, symbolic style.
Magic Prompts	I want to draw a figure showing “Knowledge is Power”, drawing of a concept art by Mark Brooks and Brad Kunkle and Craig Mullins, trending on artstation,4k
Our Prompts	I want to draw a figure showing “Knowledge is Power”. A group of ragged children huddled in a dilapidated classroom, studying intently by the faint light streaming through cracked windows, icy wind curling around them. Despite the cold and decay, a soft, ambient glow surrounds their books, casting warm light onto their faces – as if knowledge itself is illuminating them. The room is bleak and gray, but the children’s focused expressions and the golden light symbolize hope and strength. Vibrant contrast between the dull surroundings and the inner light. Cinematic composition, slightly low angle, moody and atmospheric.

Figure 10: Four Prompt Types for Constructing Images Based on the Theme 'Knowledge is Power'

Theme / By method	Love Transcends Distance
Original Prompts	I want to draw a figure showing "Love Transcends Distance".
Extended Prompts	I want to draw a figure showing "Love Transcends Distance". Zhinu and Niulang gaze at each other from afar, separated by a shimmering Milky Way. Starlight and drifting cosmic dust surround them in silence.
Magic Prompts	I want to draw a figure showing "Love Transcends Distance". Love Transcends Distance", trending on artstation, concept art.
Our Prompts	I want to draw a figure showing 'Love Transcends Distance' – two people reaching out to each other from opposite sides of a vast landscape or starry sky.

Figure 11: Four Prompt Types for Constructing Images Based on the Theme 'Love Transcends Distance'

Theme / By method	Time Heals All Wounds
Original Prompts	I want to draw a figure showing “Time Heals All Wounds”.
Extended Prompts	I want to draw a figure showing “Time Heals All Wounds” – a person sitting quietly as seasons change around them, scars on his body gradually fading into blooming flowers, symbolizing healing and renewal."
Magic Prompts	I want to draw a color figure “Time Heals All Wounds” by Marc Simonetti, trending on artstation, Midjourney.
Our Prompts	I want to draw a figure showing “Time Heals All Wounds”. A split-scene landscape showing 'Time Heals All Wounds' – on the left, charred ruins and ash-covered ground from a fire, barren trees under a gray, overcast sky; on the right, the same place fully rebuilt with tidy villas, blooming gardens, and children playing under warm sunlight. A gentle transition of light and color connects the two halves, symbolizing recovery and hope. Soft ambient lighting, vibrant greenery, emotional yet calm atmosphere, cinematic wide-angle view.

Figure 12: Four Prompt Types for Constructing Images Based on the Theme 'Time Heals All Wounds'

Theme / By method	World Peace
Original Prompts	I want to draw a figure showing the world will have peace in the end.
Extended Prompts	I want to create an image that speaks of hope and unity—a central, radiant figure standing at the heart of a softly glowing globe. This figure, draped in flowing robes of gentle blues and whites, reaches out with open arms toward a horizon where the first light of dawn touches a landscape of lush meadows, calm seas, and ancient trees. Around the figure, symbols of peace—fluttering doves, delicate olive branches, and intertwined hands—float in a serene dance. The overall atmosphere is one of quiet strength and optimism, capturing the timeless promise that no matter the challenges, the world will ultimately embrace peace.
Magic Prompts	I want to draw a figure showing the world will have peace in the concept art, low angle, high detail, warm lighting, volumetric, godrays, vivid, beautiful, trending on artstation, by Jordan grimmer, huge scene, grass, art greg Rutkowski.
Our Prompts	I want to draw a figure showing the world will have peace in the end. A breathtaking scene of global peace: a sprawling meadow ablaze with vibrant flowers, edged by a gently winding river. Diverse people—men, women, and children—play and laugh under warm, golden sunlight, while graceful birds soar in a clear blue sky, symbolizing hope and freedom.

Figure 13: Four Prompt Types for Constructing Images Based on the Theme 'World Peace'