

Personalization as Inverse Planning: Learning Latent Design Intents for Agentic Slide Generation via Structural Denoising

Tianci Liu^{1,*}, Zihan Dong², Linjun Zhang², Haoyu Wang³, Jing Gao¹,
Emre Kıcıman⁴, Ranveer Chandra⁴, and Wei-Ting Chen^{4†}

¹ Purdue University

² Rutgers University

³ University at Albany

⁴ Microsoft

Abstract. Slide design requires personalizing both deck themes and page layouts. Yet, current AI agent-based methods struggle with fine-grained, page-level design. Solely relying on prespecified templates or user verbose instructions, they fail to capture latent design intents, leaving Page-level Slide Personalization (PSP) unresolved. To close this gap, this work formulates PSP as an inverse planning problem. We propose to learn a design intent without assuming any knowledge of the specific executing tools (e.g., PowerPoint, Beamer) being used. However, relinquishing control over these tools makes the problem intractable to optimize end-to-end. To overcome this, we propose SPIRE, a principled framework to solve PSP approximately. By intentionally corrupting the visual structures of clean slides, SPIRE creates a verifiable task to denoise the corruption, whereby two agents learn to collaboratively refine executable designs via reinforcement learning (RL). We present a proof that structural denoising is a consistent surrogate for PSP, and that the multi-agent formulation strictly reduces policy gradient variance in RL. Extensive experiments demonstrate the superiority of SPIRE.

1 Introduction

Slide decks are a primary medium for communicating ideas in academia and industry [12, 33, 38]. Yet, creating a high-quality deck is a design-intensive process. Depending on the presenter and the target audience, the *same* content can call for different visual treatments. These treatments are often personalized to a speaker’s habitual design, a lab’s visual identity, or an organization’s brand guidelines. In short, practical slide generation is inherently a *personalized* task.

* Work done during an internship at Microsoft.

† Corresponding author.

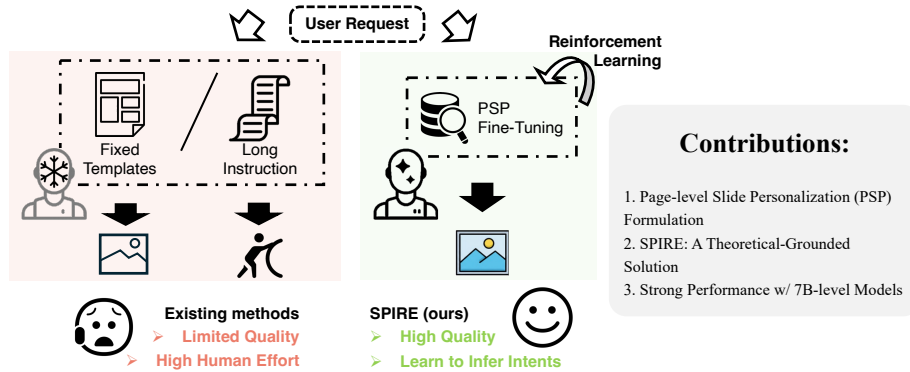


Fig. 1: Illustration of SPIRE. Trained with reinforcement learning, SPIRE learns to infer design intents instead of using prespecified layout templates or lengthy user instructions as existing methods do [9, 51], offering both *theoretical* and *empirical* advantages.

Recent advances in multi-modal large language models (MLLMs) have enabled agentic pipelines that automate slide generation for practical use [8, 9, 25, 29, 36, 43, 51]. These systems typically decompose slide creation into modular stages such as outlining, asset extraction, layout arrangement, and iterative refinement, thereby improving deck-level coherence via template selection, visual feedback, and multi-agent coordination [14, 15, 29, 40, 48, 51]. However, they struggle with fine-grained *page-level* design, which requires deciding visual hierarchy, element alignment, spacing, and styling choices. Nonetheless, existing agentic systems handle page-level layout and styling passively: Once the content is specified, the page-level design is largely overlooked, either by following generic system-defined templates [14, 23, 40, 51], or by requiring lengthy, explicit user instructions [9, 29], instead of inferring the user’s latent intent. Consequently, they are too coarse to achieve satisfactory Page-level Slide Personalization (PSP).

To close this gap, we introduce the task of agentic PSP; the concept is depicted in Fig 1. Our formulation is inspired by how human designers teach slide making in practice: *Given* a detailed plan elaborating the page specification such as layout structure and visual hierarchy, a generally capable executor (e.g., an experienced designer unfamiliar with specific corporate styles) can reliably reproduce a high-quality personalized slide visual by following the plan step-by-step.

But such actionable plans are infeasible to collect at scale in practice. Detailed intent annotations are rarely available in real-world slide corpora [9], making direct supervised training impracticable. As a result, PSP requires one to *proactively infer* the plan as a latent intent in order to actively guide the generation process. Furthermore, this plan is naturally context-dependent: (enterprise) users often maintain more than one set of intents depending on their specific scenarios. Fortunately, such contextual intent can be largely specified by having the user provide a small number of reference slides. To this end, we formulate PSP as a probabilistic problem that explicitly models the design intent as a la-

tent random variable. Given (1) user-provided page assets and (2) a small set of reference slides, we infer the final, detailed plan, which is then seamlessly passed to a downstream visual designer for execution. Intuitively, our formulated PSP aims to infer a plan that, once executed by a capable, non-personalized executor, can most reliably reproduce the desired personalized slide given the user’s references. *This probabilistic formulation treats design intent as a latent variable and provides a principled foundation for page-level slide personalization.*

Crucially, this latent-intent formulation of PSP is computationally challenging to solve for two reasons. First, it evaluates a proposed plan based on whether it enables the executor to reproduce the desired visual. This requires a rendering likelihood, which lacks a clear, direct objective that one can optimize. Naive image-level similarity provides an unreliable estimate for slide quality, which is governed by discrete, structured design decisions rather than purely pixel-wise closeness [9, 32, 37, 42]. Second, the executor that turns a plan into a slide is effectively a black box, which further impedes optimizing the PSP objective with standard numerical methods [16, 34]. Therefore, an effective approximation is needed to make this objective tractable. We detail this formulation and its computational challenges in Sec. 2.1.

As a remedy, in this work we propose SPIRE (Structural Planning via Inverse REconstruction), a structural denoising framework that turns the hard-to-optimize *slide* personalization objective into a verifiable reconstruction problem that serves as a good proxy. Our solution extends existing multi-agent visual generation frameworks [10, 13, 14, 22, 39] to a principled training objective. With the aforementioned PSP goal of learning a reference-conditioned page-level plan, we alternate between a critic that provides semantic feedback and a planner that updates its plan accordingly. The key idea is to exploit the discrete structural nature of slides: instead of perturbing pixels, we intentionally corrupt a gold slide by applying random, element-level structural perturbations to its layout, hierarchy, and styling, and record the corresponding discrepancies. *This yields scalable self-supervision where suboptimal parts to improve are known by construction.*

Built upon the structural perturbation, SPIRE trains two complementary components separately on this shared signal. A critic learns to read the corrupted slide and the user’s references, and to produce structured, actionable feedback that pinpoints the discrepancies as semantic edit suggestions. A planner learns to produce an executable, editable plan, either proposing a plan from scratch or revising an imperfect plan guided by the critic’s feedback. As both components are trained to recover from controlled structural corruptions, the critic’s feedback becomes verifiable, and the planner learns to improve plans without relying on gradients through the black-box executor, making latent-intent PSP tractable in practice. We resort to reinforcement learning [35, 45] to train the two agents. Sec. 2.2 details the proposed SPIRE. We provide theoretical analysis of its superiority in Sec. 2.3. *This principled method and its theoretical advantage are the main technical contribution of this work.*

Our paper is organized as follows. Sec. 2 formulates PSP and details the proposed SPIRE. Extensive experimental results in Sec. 3 demonstrate the effec-

tiveness of our method. In the remaining part of this paper, we review related work in Sec. 4, and conclude the paper in Sec. 5.

2 Proposed Method

This section formalizes the task of reference-based page-level slide personalization (PSP), which entails an intractable optimization problem due to contextual intent being latent. We propose SPIRE as a principled approximate solution.

2.1 Problem Formulation

Suppose a (enterprise) user specifies a high-level instruction x about slide content (e.g., text content and visual elements)⁵, and provides K reference slides

$$\mathcal{D}_{\text{ref}} = \{s_{\text{ref}}^{(1)}, \dots, s_{\text{ref}}^{(K)}\},$$

to exemplify their contextually preferred style. We define PSP as producing a target slide s that can accurately reflect the user’s exact intent that is not elaborated in the verbal x . To reflect the need for creative design, we allow $\mathcal{D}_{\text{ref}}$ not to cover the optimal layout s should use.

We refer to this intent as a *plan* and consider it as a latent variable denoted by z . Intuitively, z can be a step-by-step actionable instruction that contains a rich page-level specification (e.g., layout coordinates and element styling).

Following a well-formed z , a generally capable executor E can be unambiguously guided to produce a high quality visual s that is *personalized* to the user’s needs, even if E itself is not directly tailored to the user’s preference. From a probabilistic perspective, this implies that the visual realization s is *conditionally independent* of the user context given the plan z :

$$p(s \mid z, x, \mathcal{D}_{\text{ref}}) = p(s \mid z).$$

Here the randomness is induced by the use of black-box E , which can be a coding agent [1, 5, 6, 15, 26, 28, 50], an image generator [7, 21, 27, 44], or a human designer. Note that disentangling the plan z and E is on purpose, so that PSP does not need to be conducted whenever a new executor E is introduced.

Based on this definition, letting π_θ denote a multi-modal language model *planner*, PSP can be formulated as learning π_θ from a personal corpus

$$\mathcal{D}_{\text{tr}} = \{(x^{(1)}, (s^*)^{(1)}, \mathcal{D}_{\text{ref}}^{(1)}), \dots, (x^{(J)}, (s^*)^{(J)}, \mathcal{D}_{\text{ref}}^{(J)})\}.$$

Here s^* denotes the gold slide for x . We detail the construction of $\mathcal{D}_{\text{tr}}$ in the supplementary material.

⁵ In this paper, we refer to page-level objects on a slide (e.g., text boxes, images, shapes, charts, and tables) uniformly as *visual elements*. A *text box* is a type of visual element that contains *text content*.

Given $\mathcal{D}_{\text{tr}}$, the goal of PSP is to maximize the marginal likelihood of s^* given $x, \mathcal{D}_{\text{ref}}$. This can be expressed as the following optimization objective:

$$\begin{aligned}\mathcal{J}(\theta) &= \mathbb{E}_{(x, s^*, \mathcal{D}_{\text{ref}}) \sim \mathcal{D}_{\text{tr}}} [\log p_{\theta}(s^* | x, \mathcal{D}_{\text{ref}})] \\ &= \mathbb{E}_{(x, s^*, \mathcal{D}_{\text{ref}}) \sim \mathcal{D}_{\text{tr}}} \left[\log \int \pi_{\theta}(z | x, \mathcal{D}_{\text{ref}}) p(s^* | z) dz \right].\end{aligned}\quad (1)$$

At a colloquial level, Eq. (1) seeks π_{θ} that can generate plans that can maximize the likelihood of black-box E to reproduce (generate) the gold s^* . Unfortunately, optimizing Eq. (1) is prohibitive due to two reasons.

First, the likelihood $p(s^* | z)$ lacks an explicit form, which makes the objective intractable. One common solution is to approximately optimize

$$p(s^* | z) \propto -\|s^* - E(z)\|_2,$$

in the pixel space, aiming to push generated $s = E(z)$ towards s^* pixel-wise. Yet, recent works [9, 37, 42] showed that treating slides as pure images struggles to capture their discrete logical structure, offering unreliable guidance for the planner.

Second, black-box E cannot be differentiated through, which prevents direct computation of the gradient for π_{θ} with respect to $s = E(z)$ given by

$$\nabla_{\theta} \log p_{\theta}(s = E(z) | x, \mathcal{D}_{\text{ref}}).$$

In addition, zeroth-order optimization methods [4, 11, 18, 24] **prove** ineffective for this context for two reasons. First, their high computational cost [30, 31, 49] makes them prohibitive for slide generation, where user preferences and instructions take diverse forms. Second, these methods approximate gradients by probing the specific executor E , which makes them prone to overfitting on **a** particular instance [11], leading to limited generalizability of π_{θ} to other executors [31].

In the next section, we show that by leveraging the *structural* discreteness of slide visuals, we can circumvent these hurdles via a surrogate denoising objective, and we propose SPIRE as an effective approximate solution.

2.2 SPIRE: An Effective Solution for PSP

To tackle the computational challenge of Eq. (1), we approximate the two parts therein with two complementary agents based on their intrinsic goals. The two agents are trained to exploit the discrete structural nature of slides via a shared structural perturbation process, turning the problem into a verifiable reconstruction signal through a *denoising* process over the slide’s discrete structural space. We dub our method Structural Planning via Inverse REconstruction (SPIRE). Fig 2 depicts the overview of PSP and SPIRE.

We begin by checking the roles of Eq. (2)

$$p_{\theta}(s^* | x, \mathcal{D}_{\text{ref}}) = \int \underbrace{\pi_{\theta}(z | x, \mathcal{D}_{\text{ref}})}_{\text{planner proposal}} \underbrace{p(s^* | z)}_{\text{render likelihood}} dz, \quad (2)$$

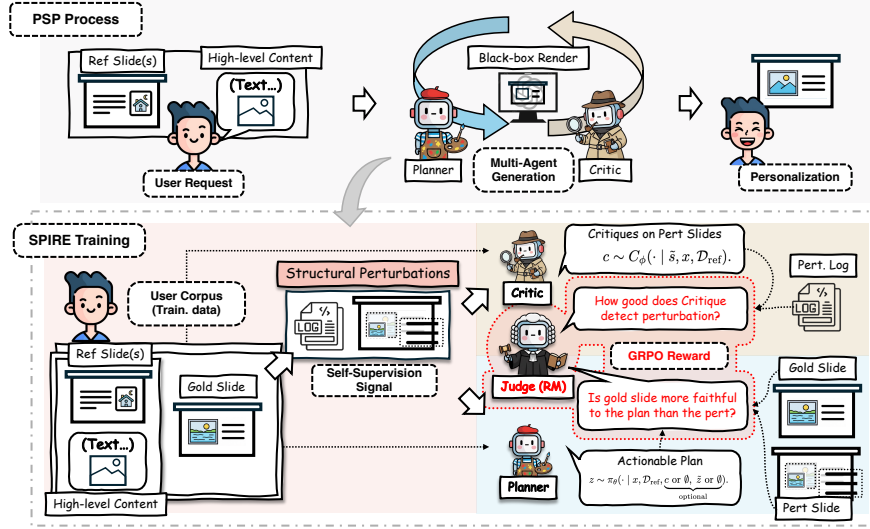


Fig. 2: The overview of Personalized Slide Personalization (PSP) and the proposed SPIRE. The Planner and Critic are trained to recover self-supervised Structural Perturbation in a decomposed way with reinforcement learning (red). After training, they collaborate with a black-box render executor to conduct PSP.

where the integrand consists of two parts. The first term, $\pi_\theta(z | x, \mathcal{D}_{\text{ref}})$, denotes how the planner π_θ proposes a plan z to reflect the user’s latent intent. Next, the second term $p(s^* | z)$ measures the *render* likelihood of the golden slide s^* , depicting the plan’s *visual* validity. PSP wants to adjust the planner proposal based on $p(s^* | z)$, which, unfortunately, is intractable due to black-box E .

To bypass this challenge, we approximate the update by incorporating a *critic* agent C_ϕ to act as a learned proxy for the likelihood. By learning to *discriminate* how a generation s visually deviates from the gold s^* , the critic captures the user’s implicit preferences and provides actionable semantic feedback, thereby providing guidance for the planner to update its proposal $\pi_\theta(z | x, \mathcal{D}_{\text{ref}})$ in context $\mathcal{D}_{\text{ref}}$. With a reliable critic, the planner can iteratively improve its proposal: starting from an initial plan, it revises the plan based on the critique.

Yet, reliably learning such a critic, and thus enabling critique-guided planner updates, is also non-trivial. The difficulty is the lack of *verifiable* supervision: for a random generation s , one can hardly evaluate whether a critique is objectively correct, which makes subsequent planner training unreliable as well.

Denoising Signal. To solve this problem, we propose a “denoising” objective through a corruption-reconstruction strategy. Namely, we construct a verifiable self-supervised signal that can be shared across the two agents’ training. Specifically, we corrupt the gold slide s^* into a perturbed slide $\tilde{s}$ in a controlled way, and record the ground-truth corruption operations. Next, the critic is asked to criticize $\tilde{s}$ based on visual evidence, and its critique can be verified by checking

the recorded corruption operations. The same corruption record also helps train the planner to perform critique-guided recovery in the *plan space*: it learns to revise a suboptimal plan $\tilde{z}$ using the critique c . We will detail this soon.

Structural Noises. We leverage the discrete structural nature of slides to design *structural corruptions*, rather than adding pixel-level noise to the visuals. Specifically, we represent each slide s^* as a collection of visual elements that incorporate different element types. Then, we apply random perturbation to each visual element by modifying its structural attributes. Each element type admits its own perturbation spaces, e.g., position for layout, size for hierarchy, and color palette for styling. The corruption occurs along three dimensions: *spatial layout* (e.g., shifting coordinates), *visual hierarchy* (e.g., distorting element sizes), and *stylistic coherence* (e.g., altering text colors). To prevent spurious correlations, we independently perturb each randomly selected element.

Formally, assume s^* contains N visual elements $\{e_i\}_{i=1}^N$, and let $\mathcal{T}(e_i)$ denote the element type (e.g., text box, image, shape). We represent the corruption by an element-level action list $\mathcal{A} = \{a_i\}_{i=1}^N$, where each a_i perturbs some attributes of e_i or leaves it unchanged. Perturbation actions are sampled independently, i.e.,

$$q(\mathcal{A} \mid s^*) = q(a_1, \dots, a_N \mid s^*) = q(a_1, \dots, a_N \mid e_1, \dots, e_N) \stackrel{(a)}{=} \prod_{i=1}^N q_{\mathcal{T}(e_i)}(a_i \mid e_i),$$

where (a) holds by independence. Here $q_{\mathcal{T}(e_i)}(\cdot \mid e_i)$ denotes a perturbation distribution over actions for element e_i (admitting no perturbation). Subscript $\mathcal{T}(e_i)$ indicates that action spaces depend on the element types. Having the sampled action list $\mathcal{A}$, we apply these actions to s^* and denote the corrupted slide by $\tilde{s}$. The inverse of each applied perturbation forms a discrepancy list:

$$\mathcal{A}_{\text{diff}} = \{a_i^{-1} \mid a_i \neq \emptyset, 1 \leq i \leq N\},$$

where $\emptyset$ denotes leaving the element unchanged. This yields self-supervised triplets $(\tilde{s}, s^*, \mathcal{A}_{\text{diff}})$, providing scalable supervision without manual annotation. Built upon the structural perturbation, SPIRE factorizes the intractable PSP problem from Eq. (1) into the following two complementary sub-tasks.

Structural Discrimination. This task aims to train critic C_ϕ to identify and articulate the discrepancies in $\mathcal{A}_{\text{diff}}$, which correspond to moving $\tilde{s}$ toward the gold s^* to maximize the *render likelihood* term in Eq. (2). Specifically, the critic generates a critique c based on x , a suboptimal visual $\tilde{s}$, and reference $\mathcal{D}_{\text{ref}}$:

$$c \sim C_\phi(\cdot \mid \tilde{s}, x, \mathcal{D}_{\text{ref}}).$$

Here, c encompasses a list of issues and targeted corrections, translating errors in $\tilde{s}$ into actionable textual feedback. We evaluate c against the full action list $\mathcal{A}$ with a capable LLM judge (GPT-4o-mini). For each element e_i ($1 \leq i \leq N$), the judge verifies whether c correctly handles that element:

$$y_i = J(c, e_i, a_i) \in \{0, 1\}, \quad 1 \leq i \leq N.$$

Here $y_i = 1$ indicates a correct decision: the critique identifies and corrects the perturbation when $a_i \neq \emptyset$, or correctly refrains from reporting an issue when $a_i = \emptyset$; and 0 otherwise. Subsequently, we define the element-level verification reward as the average verification accuracy over all N elements:

$$R_{\text{vfy}}(c, \mathcal{A}) = \frac{1}{N} \sum_{i=1}^N y_i.$$

We train the critic using DAPO [45], a strong reinforcement learning algorithm, to maximize the expected verification-format reward:

$$\mathcal{J}_{\text{critic}}(\phi) = \mathbb{E}_{\substack{s^* \sim \mathcal{D}_{\text{tr}} \\ \tilde{s} \sim q(\cdot | s^*)}} \left[\mathbb{E}_{c \sim C_\phi(\cdot | \tilde{s}, x, \mathcal{D}_{\text{ref}})} [R_{\text{vfy}}(c, \mathcal{A}) \times R_{\text{format}, c}(c)] \right]. \quad (3)$$

We detail critique generation and format requirement in the supplementary material. Note that our perturbation design, combined with element-level evaluation, mitigates reward hacking. As each perturbation type is applied independently with probability 1/2, trivial always-issue or never-issue critiques cannot consistently score high.

Structural Planning. This task aims to produce actionable plans that translate user instructions into high-quality slides. The goal is to approximately improve the *planner proposal* term in Eq. (2). To this end, we train the planner to (i) generate a good plan proposal, and (ii) update a proposal using the critic’s critique as a semantic gradient. Formally, we refer to this goal of the planner as *iterative refinement*, which covers *initial proposal* (generating z from scratch) and *critique-guided refinement* (revising a suboptimal $\tilde{z}$ based on critique c). The two subgoals can be expressed as a unified formulation. Given instruction x , the reference $\mathcal{D}_{\text{ref}}$, and optionally a suboptimal plan $\tilde{z}$ to revise, with its corresponding critique c , the planner aims to generate a high quality plan z as

$$z \sim \pi_\theta(\cdot | x, \mathcal{D}_{\text{ref}}, \underbrace{c \text{ or } \emptyset, \tilde{z} \text{ or } \emptyset}_{\text{optional}}). \quad (4)$$

The visual quality of a plan z is measured with a discriminative *plan-visual matching reward*. Specifically, a capable VLM judge (Claude Opus 4.5) is given the plan z and a pair of visuals $(s^*, \tilde{s})$, and is asked to choose which one matches the plan z better. We reward z if the judge prefers the gold slide s^* . The judgment is conducted twice with the presenting order of the two visuals swapped to mitigate the positional bias [41]. Given the presenting order $(s^*, \tilde{s})$, we denote the judgment outcome as

$$\hat{o}_{\rightarrow} = \mathbb{I}[s^* \succ \tilde{s} | z, (s^*, \tilde{s})] \in \{0, 1\},$$

and $\hat{o}_{\leftarrow}$ is defined similarly for the swapped presenting order. Based on this, we define the *plan-visual matching reward* as

$$R_{\text{pvm}}(z, s^*, \tilde{s}) = \frac{1}{2} \left(\mathbb{I}[\hat{o}_{\rightarrow} = 1] + \mathbb{I}[\hat{o}_{\leftarrow} = 1] \right)$$

To make R_{pvm} reliable, we instruct the capable judge to base its decision *solely* on fidelity to z rather than its own aesthetic preference. Note that $\tilde{s}$ is produced through structural perturbations that alter *design styles* (e.g., color palette) instead of making the slide visually worse. Without seeing x and $\mathcal{D}_{\text{ref}}$, $\tilde{s}$ may look good on its own; therefore, correct *judgment* requires a discriminative z . As with the critic, we train the planner with DAPO to maximize the expected reward, combined with a format verification term:

$$\mathcal{J}_{\text{plnr}}(\theta) = \mathbb{E}_{\substack{s^* \sim \mathcal{D}_{\text{tr}} \\ \tilde{s} \sim q(\cdot | s^*)}} \left[\mathbb{E}_{c^* \sim \text{Oracle}(c)} \left[\mathbb{E}_{z \sim \pi_\theta(\cdot | x, \mathcal{D}_{\text{ref}}, c^*)} [R_{\text{pvm}}(z, s^*, \tilde{s}) \times R_{\text{format}, \text{p}}(z)] \right] \right].$$

We defer more details about plan generation and format requirements to the supplementary material. Note that the iterative refinement goal designed for the planner requires high-quality $(c^*, \tilde{z})$ to train the agent so that it can reliably refine a given plan under critique. To this end, we use an oracle VLM (GPT-5) to (i) describe the perturbed slide $\tilde{s}$ as a suboptimal plan $\tilde{z}$, and (ii) translate the discrepancy list $\mathcal{A}_{\text{diff}}$ (provided as *hints* [17, 46, 47]) into gold critique c^* :

$$\tilde{z} \sim \text{Oracle}(\tilde{z} | \tilde{s}), \quad c^* \sim \text{Oracle}(c | \mathcal{A}_{\text{diff}}).$$

2.3 Theoretical Analysis

We end up this section with the theoretical advantages of SPIRE. Full versions are deferred to the supplementary material due to page limit.

First, under regular conditions, optimizing the SPIRE objective gives a gradient-level approximation to that of the original PSP objective up to an explicit error bound, as formalized in the following Thm 1.

Theorem 1 (Surrogate Validity for PSP, informal). *Let $\hat{\mathcal{J}}_{\text{SP}}(\theta, \phi)$ be the empirical loss of SPIRE. Then there exists $\epsilon_{\text{tot}} \geq 0$ such that*

$$\left\| \nabla_\theta \hat{\mathcal{J}}_{\text{SP}}(\theta, \phi) - \frac{\beta}{4} \nabla_\theta \mathcal{J}(\theta) \right\| \leq \epsilon_{\text{tot}}.$$

Here ϵ_{tot} aggregates critic-calibration, structural-surrogate, variational-gap, and empirical optimization errors. Full definitions are provided in the supplementary material.

Crucially, the key error term is controlled thanks to our critic training on structural corruptions with verifiable ground truth supervision. An untrained or pixel-denoising-trained C would not satisfy this bound in general.

Furthermore, Thm 2 below proves that SPIRE trains the planner and critic *separately*, thereby eliminating executor-induced noise.

Theorem 2 (Two-agent Decomposition Stabilizes Optimization, informal). *Let $\hat{g}_{e2e}$ and $\hat{g}_{2a}$ denote the end-to-end and two-agent policy-gradient estimators, respectively; see the supplementary material for explicit definitions and derivation. Then*

$$\text{Var}(\hat{g}_{e2e}) = \text{Var}(\hat{g}_{2a}) + \Delta_{\text{exec}|c}, \quad \Delta_{\text{exec}|c} \geq 0.$$

So $\text{Var}(\hat{g}_{2a}) \leq \text{Var}(\hat{g}_{e2e})$, with strict inequality for non-zero executor-induced noise. See formal statement, expression, and imperfect-critic extension in the supplementary material.

In general, even if a planner-executor-critic design is adopted, standard training of the planner and critic still requires the executor to generate visuals for complete rollouts, retaining its induced noise. SPIRE provides a way to bypass the executor in our decomposed training, thereby enabling the noise reduction.

3 Experiment

We evaluate the performance of SPIRE and ablate its components’ contributions. Benefiting from the principled optimization presented before, SPIRE offers strong PSP capability over strong baselines that rely on much larger GPT models.

3.1 Dataset and Experiment Settings

Datasets. We build the **train** and **test** data from Zenodo10k [51] by randomly sampling 200 decks. For each deck, we perform a slide-level held-out split that preserves the deck’s chronological order. The last 20% of slides are reserved for testing and the remaining slides are used for training (or for retrieval for training-free baselines). SlideBench [9] is used as **out-of-distribution** test data. We defer more data preparation details in the supplementary material.

Backbones. Both the critic and the planner are fine-tuned from Qwen2.5-VL-7B-Instruct [2]. At inference time, the planner and critic collaborate with a black-box executor E to perform multi-round iterative generation. We use GPT-o4-mini as the coding executor to implement the plan with `python-pptx`, and follow [9, 37] to improve coding reliability. No template is used in execution.

Baselines. We compare SPIRE against representative systems that cover both slide-native pipelines and generic visual synthesis. We include AutoPresent [9], which generates a slide by deciding layout and style for the given instruction from scratch on its own. We also include PPTAgent [51], which selects a template/style from the reference slides and re-applies it to the target content through an edit-based workflow. Finally, following [9], we include a generic image generation baseline using Stable Diffusion 3.5 [7] equipped with IP-Adapter [44] which prompt the model to synthesize a slide-like image. To isolate the benefit of our training, we additionally report PSP results with planner/critic setting to: (i) frontier VLM (o4-mini) and (ii) the corresponding base model (without fine-tuning), while keeping the rest of the pipeline unchanged.

Evaluations. We evaluate the generation in terms of visual similarity against the gold slides, and reference-free quality. Following [9, 20, 37], we adopt both visual-similarity metrics and a VLM-as-a-Judge protocol. Specifically, for visual similarity, we treat the slides as rendered images and utilize metrics such as SSIM and CLIP to measure “reconstruction” quality. For VLM-as-a-judge evaluation, we follow [9, 20, 52] and evaluate the generation quality from the perspectives of *faithfulness*, *color*, *layout*, and *overall aesthetics*. We detail these in the supplementary material.

Table 1: Quantitative results on *test* and *OOD* pages. Best average results are in **bold** and the second best results are underlined.

Model	Visual Similarity			VLM-as-Judge				
	Sim _{ssim} ↑	Sim _{clip} ↑	AVG	Faith ↑	Color ↑	Layout ↑	Aest ↑	AVG
Test Pages								
<i>GPT-based Models</i>								
AutoPresent [9]	0.6062	0.4431	0.5247	0.6174	0.5850	0.4145	0.4105	<u>0.5069</u>
PPTAgent [51]	0.5126	0.5491	0.5309	0.1036	0.5670	0.4364	0.3961	0.3758
PSP (o4-mini)	0.8440	0.7683	0.8062	0.5504	0.5763	0.4206	0.3661	0.4784
<i>7B-level Models</i>								
SD 3.5 [7]	0.3372	0.4496	0.3934	0.1559	0.4282	0.2545	0.1889	0.2569
PSP (Base)	0.3578	0.3317	0.3448	0.2995	0.4588	0.2956	0.2382	0.3230
SPIRE (Ours)	0.8134	0.7018	<u>0.7576</u>	0.7072	0.6330	0.4470	0.3787	0.5415
OOD Pages								
<i>GPT-based Models</i>								
AutoPresent [9]	0.5976	0.5892	0.5934	0.7923	0.7242	0.5477	0.5201	<u>0.6461</u>
PPTAgent [51]	0.4836	0.6371	0.5604	0.0839	0.5824	0.4392	0.3831	0.3722
PSP (o4-mini)	0.7233	0.7633	0.7433	0.6975	0.5947	0.4472	0.3958	0.5338
<i>7B-level Models</i>								
SD 3.5 [7]	0.3177	0.5753	0.4465	0.1811	0.5075	0.2934	0.2246	0.3016
PSP (Base)	0.3910	0.3864	0.3887	0.2432	0.2309	0.1716	0.1292	0.1937
SPIRE (Ours)	0.6785	0.6954	<u>0.6870</u>	0.9330	0.7545	0.6565	0.5891	0.7333

3.2 Quantitative Results

Tab 1 shows quantitative results on test and out-of-distribution (OOD) pages. We note the following observations.

Good PSP Performance. From the table, SPIRE achieves the strongest overall performance, outperforming the GPT-based AutoPresent (0.5069) and PSP (o4-mini) (0.4784) in terms of judge score, while maintaining a highly competitive visual similarity score (0.7414), substantially higher than the 7B-level baselines SD 3.5 and PSP (Base). This performance is notable given that SPIRE relies on only two 7B-scale agents, yet achieves competitive performance compared to GPT-based components. We also note empirical failures in the baselines, attributed to their underlying mechanisms. PPTAgent struggles with faithfulness (0.1036), as realistic reference sets rarely offer perfect templates without structural loss⁶. AutoPresent falls short when the verbose prompting is too coarse. *These trends clearly support our PSP formulation.* The suboptimal results of the training-free PSP baselines highlight the necessity of preference-aligned optimization; without it, the critic leverages its own generic preferences rather than reflecting the user’s contextual needs. *This confirms that SPIRE’s gains stem from the RL training, not merely the multi-agent refinement loop.*

⁶ In reality, `pptx` template files may not be accessible.

	Ref	AutoPre.	PPTAgent	SD 3.5	PSP (o4-mini)	PSP (Base)	Ours	Gold
Test								
								
								
OOD								
								
								

Fig. 3: Qualitative comparison of slide generation results.

Strong Generalizability. On OOD pages, SPIRE demonstrates strong generalizability and achieves the highest judge average, substantially outperforming baselines such as AutoPresent, PSP (o4-mini), PPTAgent, SD 3.5, and PSP (Base). Our method obtains the best scores across all judging dimensions and the second-best visual similarity. These results indicate that SPIRE does not simply memorize deck-specific patterns from the training corpus, but instead successfully leverages the reference slides as contextual evidence at inference time, allowing it to infer the appropriate design intent in unseen scenarios.

Metric Discrepancy. We also note that visual similarity and judge-based scores are not perfectly aligned. For instance, while PSP (o4-mini) achieves the highest visual averages across both in-distribution and OOD settings, its judge-based quality remains substantially lower than that of SPIRE. This mismatch, as explained in Sec. 2, highlights the challenge of delivering PSP by optimizing standardized metrics alone—a limitation also noted in the literature [9, 32, 42].

These results collectively demonstrate the effectiveness of the proposed SPIRE.

3.3 Qualitative Results

We provide a visual comparison of the generated slides across different methods. Illustrative results are shown in Fig 3, with more examples deferred to the supplementary material. We note that SPIRE consistently produces visually appealing and structurally coherent slides that closely align with the gold slides, while baselines often struggle to maintain visual fidelity or fail to generate valid structures.

Table 2: Ablation study on SPIRE Training and Iterative Revision. Both components are necessary for high-quality generation.

Model	Visual Similarity			VLM-as-Judge				
	Sim _{ssim} ↑	Sim _{clip} ↑	AVG	Faith ↑	Color ↑	Layout ↑	Aest ↑	AVG
w/o FT	0.3578	0.3317	0.3448	0.2995	0.4588	0.2956	0.2382	0.3230
w/o Revision	0.7299	0.6096	0.6698	0.4865	0.5120	0.3386	0.2779	0.4038
SPIRE	0.8134	0.7018	0.7576	0.7072	0.6330	0.4470	0.3787	0.5415

Table 3: Ablation study on the critic role. Substituting the larger, untrained GPT critic with SPIRE-trained 7B model results in improved performance.

Model	Visual Similarity			VLM-as-Judge				
	Sim _{ssim} ↑	Sim _{clip} ↑	AVG	Faith ↑	Color ↑	Layout ↑	Aest ↑	AVG
w/ o4-mini Critic	0.8440	0.7683	0.8062	0.5504	0.5763	0.4206	0.3661	0.4784
w/ SPIRE Critic	0.9367	0.8633	0.9000	0.5777	0.5905	0.4515	0.3942	0.5035

The effectiveness of SPIRE, stemming from our preference-aligned RL training, empowers the model to explicitly infer latent user intent. This capability manifests in two aspects. First, SPIRE achieves superior color coherence. For instance, as shown in Row 4 and 5, SPIRE accurately *infers* a proper coloring strategy that differs from the reference slide in OOD scenarios. Second, SPIRE demonstrates strong spatial reasoning for layout arrangement. In Row 2, it successfully organizes complex, multi-image visual assets (e.g., heatmaps) into a clean, aligned structure that matches the gold slide perfectly. In Row 3, it also correctly positions the logo and text blocks. This proactive inference capability allows SPIRE to surpass baselines solely relying on verbose instructions (AutoPresent) or generic templates (PPTAgent). When such explicit inputs are unavailable, these methods inevitably fail to capture the user’s true intent.

3.4 Ablation Study

We end up this section with ablation study on key components in SPIRE.

SPIRE Training and Iterative Revision. To break down the contributions of our proposed framework, we ablate the fine-tuning process and the multi-turn revision mechanism, with results reported in Tab. 2. First, we observe that preference-aligned fine-tuning is necessary for PSP. Namely, when directly prompting base Qwen models to conduct PSP, they struggle significantly and achieve a visual average of only 0.3344 and a judge average of 0.3230. In contrast, our fine-tuned SPIRE achieves much better scores, reaching 0.7414 and 0.5415 respectively. This massive performance gap confirms that off-the-shelf VLMs, especially small-scale models, cannot reliably infer latent page-level design intents without our targeted structural denoising optimization. Second, the results validate the effectiveness of the iterative revision process. As shown in the table,

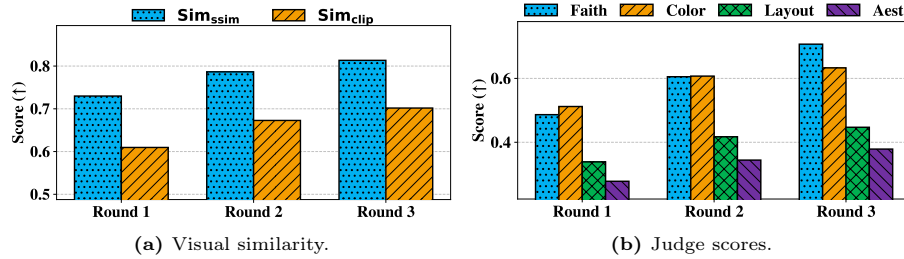


Fig. 4: Ablation on SPIRE revision. Multi-round revision in SPIRE helps improve both visual similarity and judge scores, indicating better generation qualities.

and Fig. 4, the visual similarity scores exhibit a clear and steady improvement through the successive refinement rounds. This steady climb is a direct result of our structural denoising objective, which explicitly trains the planner to interpret and execute structural corrections based on the critic’s feedback. Consequently, the planner is capable of progressively refining suboptimal layouts, proving that SPIRE successfully leverages high-quality critiques to iteratively enhance the design rather than relying on a single-pass generation.

Critic Role. One hypothesis stated in Sec. 3.2 is that the suboptimal performance of PSP lies in its critic. To verify that the preference-aligned training helps the critic learn the user’s contextual needs better, we replace the critic in the PSP (o4-mini) baseline with our 7B-level SPIRE-trained critic and report the results on test slides. Results are reported in Tab 3. From the table, we see that this substitution yields consistent improvements across all metrics. Remarkably, the critic is a 7B-level model, which is much less capable than GPT models. This gain confirms that such a targeted critic indeed helps guide the PSP process.

4 Related Work

Early works consider automated slide generation as a text summarization task [3, 8, 36], aiming to extract [8, 36] and summarize [3] proper deck-level content from a given document for presentation. These solutions boiled down to extracting salient sentences, figures, and sections from source documents to organize them into slide outlines. However, they overlooked the necessity of coherent layout design and the inherent multimodal nature of slides [19, 48]. As a result, these solutions offered limited control over the aesthetic perspective of the generated slides, leaving personalized generation untouched.

Following the emergence of (M)LLMs, recent research has shifted toward end-to-end and agentic pipelines for slide design [9, 14, 15, 25, 29, 40, 48, 51]. Representative solutions decompose the task into modular stages, including content outlining, asset extraction, layout arrangement, and iterative refinement. These systems improve visual consistency through template-based selection, visual feedback, and multi-agent coordination [15, 25, 29, 40, 48, 51]. Nevertheless,

these solutions mainly focus on more global deck- or document-level decisions and planning, lacking more fine-grained page-level layout design capabilities. For such page-level design, these solutions are either guided by generic templates and system-defined objectives [48, 51], or by explicit and lengthy instructions specified by users [9, 29]. These designs inherently limit their performance for PSP, which requires deep understanding of user-specific design intents.

Very recent works have explored personalization for slide generation. However, existing works mostly focus on adapting content and narrative structure to different audiences or presentation contexts [15, 48, 51]. [25] enables more tailored communication by conditioning on audience profiles or example pairs, and [48] allows users to specify a document–slide deck pair to express their preference, with the system aiming to replicate reference slides while plugging in content from the user’s own document. However, these designs still lack the ability to create novel page-level designs tailored to the user’s intent. In this work, we propose SPIRE to learn a user’s page-level preferences from their visual data, pushing agentic slide generation to a more fine-grained page-level design task.

5 Conclusion

This paper studies Page-level Slide Personalization (PSP), a key aspect of agentic slide generation that remains largely underexplored. Broadly, we show that personalized visual generation is fundamentally a latent-intent inference problem. We formulate PSP as an inverse planning problem to infer a user’s latent design intent, which can be used to guide diverse executors to render the visuals. We propose SPIRE, which leverages structural denoising to construct a tractable surrogate objective for optimization. By intentionally corrupting the visual structures of gold slides, SPIRE creates a self-supervised denoising task whereby two agents are trained to iteratively refine the design plan. We provide theoretical analysis on how structural denoising serves as a consistent surrogate objective for PSP, and how our multi-agent formulation strictly reduces policy gradient variance to stabilize the optimization. Extensive experiments demonstrate the effectiveness of our method against representative baselines for the PSP task.

Acknowledgments

This work is supported in part by the US National Science Foundation under grant NSF IIS-2141037. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the author(s) and do not necessarily reflect the views of the National Science Foundation.

References

1. Anthropic: Claude sonnet 4: Hybrid reasoning model with superior intelligence for high-volume use cases, and 200k context window. <https://www.anthropic.com/claude/sonnet> (2025)
2. Bai, S., Chen, K., Liu, X., Wang, J., Ge, W., Song, S., Dang, K., Wang, P., Wang, S., Tang, J., Zhong, H., Zhu, Y., Yang, M., Li, Z., Wan, J., Wang, P., Ding, W., Fu, Z., Xu, Y., Ye, J., Zhang, X., Xie, T., Cheng, Z., Zhang, H., Yang, Z., Xu, H., Lin, J.: Qwen2.5-vl technical report (2025)
3. Cao, J., Zhang, X., Li, R., Wei, J., Li, C., Joty, S., Carenini, G.: Multi2: Multi-agent test-time scalable framework for multi-document processing. In: Proceedings of The 5th New Frontiers in Summarization Workshop. pp. 135–156 (2025)
4. Chen, L., Chen, J., Goldstein, T., Huang, H., Zhou, T.: Instructzero: Efficient instruction optimization for black-box large language models. arXiv preprint arXiv:2306.03082 (2023)
5. Comanici, G., Bieber, E., Schaeckermann, M., Pasupat, I., Sachdeva, N., Dhillon, I., Blistein, M., Ram, O., Zhang, D., Rosen, E., et al.: Gemini 2.5: Pushing the frontier with advanced reasoning, multimodality, long context, and next generation agentic capabilities. arXiv preprint arXiv:2507.06261 (2025)
6. Dong, Y., Jiang, X., Qian, J., Wang, T., Zhang, K., Jin, Z., Li, G.: A survey on code generation with llm-based agents. arXiv preprint arXiv:2508.00083 (2025)
7. Esser, P., Kulal, S., Blattmann, A., Entezari, R., Müller, J., Saini, H., Levi, Y., Lorenz, D., Sauer, A., Boesel, F., et al.: Scaling rectified flow transformers for high-resolution image synthesis. In: Forty-first international conference on machine learning (2024)
8. Fu, T.J., Wang, W.Y., McDuff, D., Song, Y.: Doc2ppt: Automatic presentation slides generation from scientific documents. In: Proceedings of the AAAI Conference on Artificial Intelligence. vol. 36, pp. 634–642 (2022)
9. Ge, J., Wang, Z.Z., Zhou, X., Peng, Y.H., Subramanian, S., Tan, Q., Sap, M., Suhr, A., Fried, D., Neubig, G., et al.: Autopresent: Designing structured visuals from scratch. In: Proceedings of the Computer Vision and Pattern Recognition Conference. pp. 2902–2911 (2025)
10. Hahn, M., Zeng, W., Kannen, N., Galt, R., Badola, K., Kim, B., Wang, Z.: Proactive agents for multi-turn text-to-image generation under uncertainty. arXiv preprint arXiv:2412.06771 (2024)
11. Hu, W., Shu, Y., Yu, Z., Wu, Z., Lin, X., Dai, Z., Ng, S.K., Low, B.K.H.: Localized zeroth-order prompt optimization. *Advances in Neural Information Processing Systems* **37**, 86309–86345 (2024)
12. Hu, Y., Wan, X.: Ppsgen: Learning to generate presentation slides for academic papers. In: IJCAI. pp. 2099–2105 (2013)
13. Jaiswal, S., Prabhudesai, M., Bhardwaj, N., Qin, Z., Zadeh, A., Li, C., Fragkiadaki, K., Pathak, D.: Iterative refinement improves compositional image generation. arXiv preprint arXiv:2601.15286 (2026)
14. Jang, D., Heisler, M.L., Xing, L., Li, Y., Wang, E., Xiong, Y., Zhang, Y., Fan, Z.: Deckbench: Benchmarking multi-agent frameworks for academic slide generation and editing. arXiv preprint arXiv:2602.13318 (2026)
15. Jung, K., Cho, H., Yun, J., Yang, S., Jang, J., Choo, J.: Talk to your slides: Language-driven agents for efficient slide editing. arXiv preprint arXiv:2505.11604 (2025)

16. Lan, G.: First-order and stochastic optimization methods for machine learning, vol. 1. Springer (2020)
17. Li, C., Xue, M., Zhang, Z., Yang, J., Zhang, B., Yu, B., Hui, B., Lin, J., Wang, X., Liu, D.: Start: Self-taught reasoner with tools. In: Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing. pp. 13523–13564 (2025)
18. Li, W., Wang, X., Li, W., Jin, B.: A survey of automatic prompt engineering: An optimization perspective. arXiv preprint arXiv:2502.11560 (2025)
19. Liang, X., Zhang, X., Xu, Y., Sun, S., You, C.: Slidegen: Collaborative multimodal agents for scientific slide generation. arXiv preprint arXiv:2512.04529 (2025)
20. Liu, C., Yang, Y., Zhou, K., Zhang, Z., Fan, Y., Xie, Y., Qi, P., Wang, X.E.: Presenting a paper is an art: Self-improvement aesthetic agents for academic presentations. arXiv preprint arXiv:2510.05571 (2025)
21. Ma, J., Liang, J., Chen, C., Lu, H.: Subject-diffusion: Open domain personalized text-to-image generation without test-time fine-tuning. In: ACM SIGGRAPH 2024 Conference Papers. pp. 1–12 (2024)
22. Ma, S., Guo, Y., Su, J., Huang, Q., Zhou, Z., Wang, Y.: Talk2image: A multi-agent system for multi-turn image generation and editing. arXiv preprint arXiv:2508.06916 (2025)
23. Maheshwari, H., Bandyopadhyay, S., Garimella, A., Natarajan, A.: Presentations are not always linear! gnn meets llm for document-to-presentation transformation with attribution. arXiv preprint arXiv:2405.13095 (2024)
24. Malladi, S., Gao, T., Nichani, E., Damian, A., Lee, J.D., Chen, D., Arora, S.: Fine-tuning language models with just forward passes. *Advances in Neural Information Processing Systems* **36**, 53038–53075 (2023)
25. Mondal, I., Shwetha, S., Natarajan, A., Garimella, A., Bandyopadhyay, S., Boyd-Graber, J.: Presentations by the humans and for the humans: Harnessing llms for generating persona-aware slides from documents. In: Proceedings of the 18th Conference of the European Chapter of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 2664–2684 (2024)
26. OpenAI: Introducing chatgpt (2022), <https://openai.com/blog/chatgpt>
27. OpenAI: Gpt-image-1. <https://platform.openai.com/docs/models/gpt-image-1> (2025)
28. OpenAI: Introducing gpt-5. <https://openai.com/index/introducing-gpt-5/> (2025)
29. Pang, W., Lin, K.Q., Jian, X., He, X., Torr, P.: Paper2poster: Towards multimodal poster automation from scientific papers. arXiv preprint arXiv:2505.21497 (2025)
30. Park, S., Jeong, J., Kim, Y., Lee, J., Lee, N.: Zip: An efficient zeroth-order prompt tuning for black-box vision-language models. arXiv preprint arXiv:2504.06838 (2025)
31. Qi, Y., Tian, J., Liu, T., Li, R., Wei, T., Liu, H., Tang, X., Cheng, M., He, J.: Learning to instruct: Fine-tuning a task-aware instruction optimizer for black-box llms. In: Findings of the Association for Computational Linguistics: EMNLP 2025. pp. 7707–7733 (2025)
32. Qu, Y., Fang, S., Wang, Y., Wang, X., Chen, Z., Xie, H., Zhang, Y.: Igd: Instructional graphic design with multimodal layer generation. In: Proceedings of the IEEE/CVF International Conference on Computer Vision. pp. 18218–18228 (2025)
33. Rojo, M.G., García, G.B., Mateos, C.P., García, J.G., Vicente, M.C.: Critical comparison of 31 commercially available digital slide systems in pathology. *International journal of surgical pathology* **14**(4), 285–305 (2006)

34. Ruder, S.: An overview of gradient descent optimization algorithms. arXiv preprint arXiv:1609.04747 (2016)
35. Shao, Z., Wang, P., Zhu, Q., Xu, R., Song, J., Bi, X., Zhang, H., Zhang, M., Li, Y., Wu, Y., et al.: Deepseekmath: Pushing the limits of mathematical reasoning in open language models. arXiv preprint arXiv:2402.03300 (2024)
36. Sun, E., Hou, Y., Wang, D., Zhang, Y., Wang, N.X.: D2s: Document-to-slide generation via query-based text summarization. In: Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies. pp. 1405–1418 (2021)
37. Tang, W., Xiao, J., Jiang, W., Xiao, X., Wang, Y., Tang, X., Li, Q., Ma, Y., Liu, J., Tang, S., et al.: Slidecoder: Layout-aware rag-enhanced hierarchical slide generation from design. In: Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing. pp. 9026–9050 (2025)
38. Trelease, R.B.: From chalkboard, slides, and to e-learning: How computing technologies have transformed anatomical sciences education. *Anatomical sciences education* **9**(6), 583–602 (2016)
39. Venkatesh, K., Dunlop, C., Yanardag, P.: Crea: A collaborative multi-agent framework for creative image editing and generation. arXiv preprint arXiv:2504.05306 (2025)
40. Xie, E., Waterfield, D., Kennedy, M., Zhang, A.: Slidebot: A multi-agent framework for generating informative, reliable, multi-modal presentations. arXiv preprint arXiv:2511.09804 (2025)
41. Xu, R., Liu, T., Dong, Z., You, T., Hong, I., Yang, C., Zhang, L., Zhao, T., Wang, H.: Alternating reinforcement learning for rubric-based reward modeling in non-verifiable llm post-training. arXiv preprint arXiv:2602.01511 (2026)
42. Xu, X., Xu, X., Chen, S., Chen, H., Zhang, F., Chen, Y.C.: Pregenie: An agentic framework for high-quality visual presentation generation. arXiv preprint arXiv:2505.21660 (2025)
43. Xu, Y., Ma, X., Qiu, J., Zhao, H.: Textual-to-visual iterative self-verification for slide generation. arXiv preprint arXiv:2502.15412 (2025)
44. Ye, H., Zhang, J., Liu, S., Han, X., Yang, W.: Ip-adapter: Text compatible image prompt adapter for text-to-image diffusion models. arXiv preprint arXiv:2308.06721 (2023)
45. Yu, Q., Zhang, Z., Zhu, R., Yuan, Y., Zuo, X., Yue, Y., Dai, W., Fan, T., Liu, G., Liu, L., et al.: Dapo: An open-source llm reinforcement learning system at scale. arXiv preprint arXiv:2503.14476 (2025)
46. Zelikman, E., Harik, G., Shao, Y., Jayasiri, V., Haber, N., Goodman, N.D.: Quietstar: Language models can teach themselves to think before speaking. arXiv preprint arXiv:2403.09629 (2024)
47. Zelikman, E., Wu, Y., Mu, J., Goodman, N.: Star: Bootstrapping reasoning with reasoning. *Advances in Neural Information Processing Systems* **35**, 15476–15488 (2022)
48. Zeng, W., Ouyang, M., Cui, L., Ng, H.T.: Slidetaylor: Personalized presentation slide generation for scientific papers. arXiv preprint arXiv:2512.20292 (2025)
49. Zhan, H., Chen, C., Ding, T., Li, Z., Sun, R.: Unlocking black-box prompt tuning efficiency via zeroth-order optimization. In: Findings of the Association for Computational Linguistics: EMNLP 2024. pp. 14825–14838 (2024)
50. Zhang, K., Li, J., Li, G., Shi, X., Jin, Z.: Codeagent: Enhancing code generation with tool-integrated agent systems for real-world repo-level coding challenges. In: Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers). pp. 13643–13658 (2024)

51. Zheng, H., Guan, X., Kong, H., Zhang, W., Zheng, J., Zhou, W., Lin, H., Lu, Y., Han, X., Sun, L.: Pptagent: Generating and evaluating presentations beyond text-to-slides. In: Proceedings of the 2025 Conference on Empirical Methods in Natural Language Processing. pp. 14413–14429 (2025)
52. Zhu, D., Meng, R., Song, Y., Wei, X., Li, S., Pfister, T., Yoon, J.: Paperbanana: Automating academic illustration for ai scientists. arXiv preprint arXiv:2601.23265 (2026)