

M3DLayout: A Multi-Source Dataset of 3D Indoor Layouts and Structured Descriptions for 3D Generation

Yiheng Zhang^{1*} Zhuojiang Cai^{2*} Mingdao Wang^{1*} Meitong Guo¹
Tianxiao Li¹ Li Lin³ Yuwang Wang^{1†}

¹Tsinghua University ²Beihang University ³Migu Beijing Research Institute

<https://graphic-kiliani.github.io/M3DLayout/>

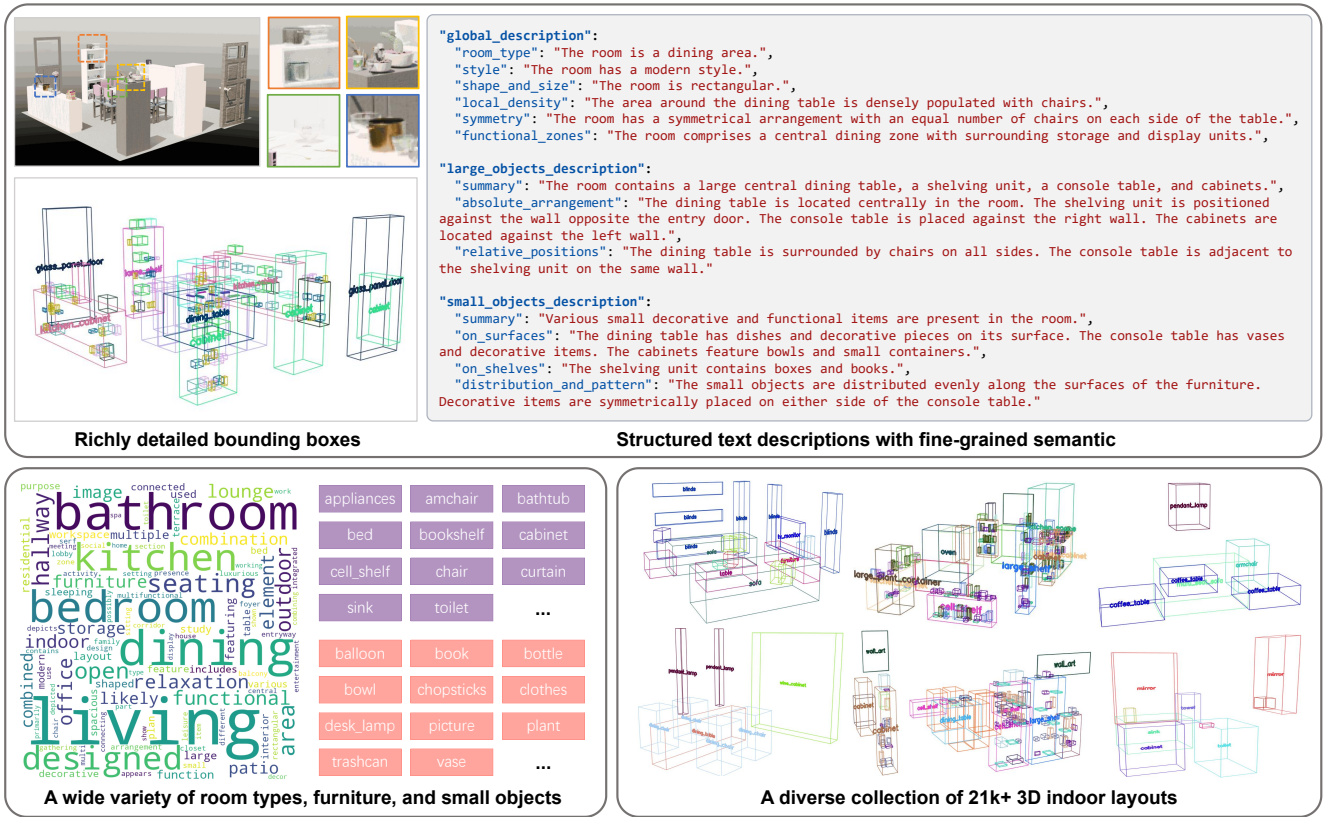


Figure 1. **The M3DLayout dataset — A multi-source benchmark for text-to-3D indoor scene generation.** Top: An example from our dataset showing a detailed 3D indoor layout with richly annotated bounding boxes and its corresponding structured textual description. Bottom-left: Word cloud visualization demonstrating the diversity of room types, furniture, and objects in the dataset. Bottom-right: Overview of the large-scale collection containing 21,367 diverse 3D layout scenes with various styles.

Abstract

In text-driven 3D scene generation, object layout serves as a crucial intermediate representation that bridges high-level language instructions with detailed geometric output. It

not only provides a structural blueprint for ensuring physical plausibility but also supports semantic controllability and interactive editing. However, the learning capabilities of current 3D indoor layout generation models are constrained by the limited scale, diversity, and annotation quality of existing datasets. To address this, we introduce **M3DLayout**, a large-scale, multi-source dataset for 3D in-

*Equal contribution

†Corresponding author

door layout generation. M3DLayout comprises 21,367 layouts and over 433k object instances, integrating three distinct sources: real-world scans, professional CAD designs, and procedurally generated scenes. Each layout is paired with detailed structured text describing global scene summaries, relational placements of large furniture, and fine-grained arrangements of smaller items. This diverse and richly annotated resource enables models to learn complex spatial and semantic patterns across a wide variety of indoor environments. To assess the potential of M3DLayout, we establish a benchmark using both a text-conditioned diffusion model and a text-conditioned autoregressive model. Experimental results demonstrate that our dataset provides a solid foundation for training layout generation models. Its multi-source composition enhances diversity, notably through the Inf3DLayout subset which provides rich small-object information, enabling the generation of more complex and detailed scenes. We hope that M3DLayout can serve as a valuable resource for advancing research in text-driven 3D scene synthesis. All dataset and code will be made public upon acceptance.

1. Introduction

Recent advances in 3D generative modeling have enabled remarkable progress in synthesizing 3D objects and scenes from various modalities, such as text or images. These developments have great potential for downstream applications in areas such as content creation, robotics, and virtual reality [14, 25, 35, 36]. In particular, text-to-3D generation has attracted increasing attention due to its intuitive and flexible interface for controlling complex scene content. For example, LucidDreamer [5] and Text2Immersion [20] adopt point-based or Gaussian-splatting representations to generate detailed scene geometry directly from text or other modalities. Other approaches, such as ATISS [21], SceneFormer [30], EchoScene [39], and MIDI [16], perform joint layout-object generation by autoregressively predicting room structures and furnishing objects in a unified framework. LayoutGPT [8], HoloDeck [36], and InstructScene [18] employ LLMs to plan scene layouts from free-form descriptions, showcasing a promising direction in LLM-driven compositional layout generation.

While these methods demonstrate strong capabilities, they also reveal key limitations. Some of them generate scenes as inseparable volumetric representations, which limits modularity and downstream controllability. Layout-aware models such as CommonScenes [38] and DiffuScene [29] tend to produce relatively simple scenes with few object types and limited diversity. Meanwhile, LLM-based planners show promise in parsing natural language but often struggle with spatial consistency and accurate physical arrangements.

To overcome this bottleneck, we draw inspiration from AutoPartGen [4], which decomposes a scene into distinct parts and thus establishes a bridge between part-level object generation and scene generation. However, AutoPartGen is restricted to small and relatively simple scenes since they simply operate under masked-image conditioning, making it difficult to scale to complex 3D scene environments. In contrast, mainstream object-level part generation methods such as OmniPart [37], X-Part [34], and CoPart [26] rely heavily on 3D bounding boxes as context to control the locations and numbers of parts, highlighting the crucial role of 3D bounding boxes as an intermediate representation. When these 3D bounding boxes are further enriched with semantic labels, rotation information and text descriptions, they form a full 3D layout that not only defines the structural backbone of the scene, but also provides powerful conditional guidance and reflects the functional intent of the environment. Building on this insight, we integrate such layout representations into the scene generation pipeline to produce high-quality, controllable 3D scenes with coherent spatial structure and realistic functionality and further unify object generation and scene generation.

To be more specific, the importance of 3D layouts can be summarized in three key points: (i) 3D layout data define the position, orientation, and scale of objects within a scene, forming the foundation of its structure and spatial coherence. This ensures objects are placed logically and functionally, greatly enhancing the generated scene’s realism and credibility while preventing chaotic or illogical visual outcomes. (ii) As powerful conditional information, 3D layout data guides and constrains the 3D scene generation process. It significantly reduces the model’s degrees of freedom and ambiguity, allowing for more efficient and accurate detail filling. This also enables users to precisely control the scene’s layout, meeting personalized and diverse generation demands. (iii) Real-world 3D scenes typically serve specific functions, and 3D layout data directly reflect this functionality. Through logical arrangements, generated scenes can better fulfill practical application needs (e.g., interior design, game levels) and provide interactive spaces that align with human habits, thereby significantly improving the end-user experience.

We argue that a major constraint limiting further progress in controllable and high-quality scene generation is the lack of large-scale, richly annotated 3D layout datasets that provide structured, semantic-level supervision. Existing datasets either focus on scene geometry from real-world scans (e.g., ScanNet [6], Matterport3D [3]) or offer object-level annotations based on professional designs (e.g., 3D-FRONT [9], Structured3D [41]). However, none of them offer comprehensive layout annotations—covering large furniture, small objects, and decorative items—paired with structured descriptions of global scene organization

and fine-grained spatial relations, and scene feasibility issues, such as overlapping, displacement, and semantic violations, are rarely verified, leaving substantial noise that ultimately limits text-to-layout generation and scene generation.

To address this gap, we introduce M3DLayout, a multi-source dataset for 3D indoor layout generation from structured language. M3DLayout contains 21,367 layouts and over 433k object instances, collected from three complementary sources: real-world scans, professional CAD designs, and procedurally generated scenes. We pair each layout with structured text descriptions that capture global scene organization, relational placements of large furniture, and fine-grained arrangements of smaller items. These high-quality annotations were constructed using a combination of rule-based extraction, GPT-assisted generation, and rigorous human verification.

To demonstrate the utility of M3DLayout, we establish a benchmark using both a text-conditioned diffusion model and a text-conditioned autoregressive model as baselines for layout generation. Experimental results show that our dataset provides a solid foundation for training layout generation models. Its multi-source composition enhances diversity, notably through the Inf3DLayout subset which provides rich small-object information, enabling the generation of more complex and detailed scenes.

We summarize our main contributions as follows:

- We introduce M3DLayout, a large-scale, richly annotated dataset of 3D indoor layouts with structured text descriptions, compiled from diverse complementary sources.
- We dedicate efforts to creating the Inf3DLayout subset, which fills a gap in high-quality data for common scene types and substantially increases the level of detail and diversity within the dataset.
- We establish two baselines for text-to-layout generation using both diffusion-based and autoregressive baselines. Results show that M3DLayout enhances the diversity and detail of generated scenes, and also enables explicit control over the fineness and density of the generated 3D layout via textual prompts.

2. Related Work

2.1. Datasets for Indoor Scenes

Large-scale datasets play a crucial role in learning-based 3D indoor layout generation and scene synthesis. Early efforts to capture real-world environments using 3D scans led to the creation of datasets such as ScanNet [6], Matterport3D [3], and SceneNN [15]. These datasets provide high-fidelity mesh reconstructions, reflect real-world object distributions and spatial constraints. However, they often suffer from noisy geometry, incomplete object coverage, and a lack of fine-grained annotations suitable for genera-

tive tasks, as well as limited layout variability due to constrained capture environments.

To address these limitations, synthetic datasets such as SUNCG [27], 3D-FRONT [9], and Structured3D [41] were introduced, offering structured 3D layouts with complete object metadata and annotations from professional CAD designs. However, these professional designs typically lack object variety and fine-grained detail.

Recent hybrid datasets attempt to bridge this gap. FurniScene [40] enriches layout realism with more diverse furniture arrangements. OpenRooms [17] provides photorealistic rendering with physical material properties. However, a key limitation persists across nearly all prior datasets: the lack of scene-level textual annotations, which limits their use for conditional or multimodal generation tasks.

In this context, we introduce M3DLayout, a multi-source dataset that combines real-world scans, professional designs, and procedurally generated layouts. We further enrich these sources with structured textual annotations, creating a foundational resource aimed at propelling research in controllable, text-driven 3D scene synthesis.

2.2. Indoor Scene Synthesis

The evolution of indoor scene synthesis methods highlights the need for richer, more descriptive layout data. Procedural approaches generate indoor scenes using predefined rules, templates, or simulation engines. Systems such as Procthor [7] and Infinigen [23, 24] rely on large-scale procedural scene grammars and asset libraries to create diverse, physically realistic environments. These methods offer high controllability and scalability, especially for generating synthetic training data for embodied agents. However, they are ultimately constrained by their handcrafted rules.

In parallel, a large body of work uses learning-based generative models trained on indoor scene datasets to learn object layouts and spatial relationships in a data-driven fashion. Early methods relied on auto-encoding architectures (e.g., SG-VAE [22], SceneHGN [10], CommonScenes [38]) and autoregressive models (e.g., ATISS [21], SceneFormer [30]). More recently, diffusion-based models such as DiffuScene [29] and EchoScene [39] have shown superior performance in capturing complex spatial dependencies. SemlayoutDiff [28] optimized the DiffuScene [29] by using semantic-mediated 2D distribution maps to finally generate 3D layouts, and is capable of generating different room types with the same model. However, precise control of small object generation cannot be achieved through functional zoning alone. These methods improve the plausibility and diversity of generated layouts but often face challenges in optimization, attribute disentanglement, and generalization to out-of-distribution scenes.

Inspired by the success of large language models (LLMs), several works explore text-guided 3D scene gen-

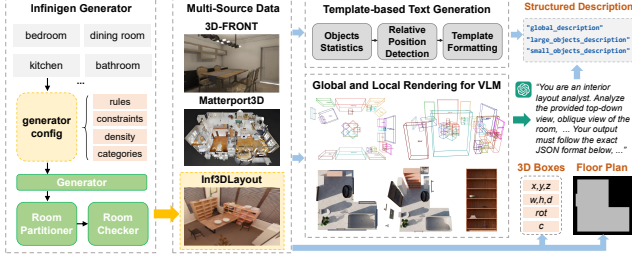


Figure 2. **Pipeline for Constructing the M3DLayout Dataset.** Our framework integrates multi-source data, including the professional designs dataset 3D-FRONT, real-world scans from Matterport3D, and procedurally generated scenes from Infinigen. The construction process involves: meticulously generating, partitioning, and filtering layouts to create the Inf3DLayout subset; performing template-based rules to produce formatted text; and employing global and local rendering for vision-language models (VLM) to produce structured descriptions. This pipeline results in a large-scale, richly-annotated text-3D layout paired dataset.

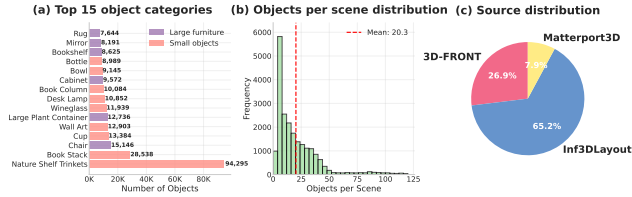


Figure 3. **Dataset statistics of M3DLayout.** (a) Top 15 most frequent object categories. (b) Distribution of the number of objects per scene. (c) Proportion of scenes contributed by each source.

eration. Architect [31] and HoloDeck [36] leverage world knowledge to interpret user instructions and generate spatial constraints. Similarly, FlairGPT [19] controls layout generation through more detailed object descriptions using LLMs, but the complex conversational process greatly reduces the efficiency of generating large amounts of data. While these methods support flexible and intuitive interactions, they often struggle with precise 3D spatial reasoning, often producing physically implausible layouts. This core limitation stems from their lack of grounding in large-scale, physically plausible 3D scene data. M3DLayout is designed to address this critical gap, providing the missing link and a robust foundation to train and benchmark the next generation of these powerful synthesis models.

3. The M3DLayout Dataset

To advance controllable and generalizable text-to-3D scene generation, we introduce M3DLayout, a large-scale, multi-source dataset of 3D indoor layouts. This dataset integrates 21,367 layouts from three complementary types of sources: real-world scans, professional interior designs, and procedurally generated scenes. Each layout is annotated

with detailed structured textual descriptions to support fine-grained, text-conditioned layout generation.

3.1. Data Sources and Curation

M3DLayout is built upon three distinct types of data sources, each contributing unique characteristics:

Real-world Scans. We incorporate layouts from the Matterport3D dataset [3], which are derived from real environment scans. These layouts reflect realistic, and often cluttered, spatial arrangements found in actual indoor settings. To ensure data quality, we performed a cleanup of the object category list by merging and removing low-frequency categories. Scenes containing fewer than two object instances were filtered out. This curated subset provides a wide variety of scene types and offers important cues for learning robust and realistic layout patterns.

Professional Interior Designs. We integrate high-quality layouts from 3D-FRONT [9], which contains professionally designed indoor scenes. These layouts are characterized by well-organized spatial semantics and adhere to tidy, minimalist design principles. They typically feature sparser object arrangements, providing strong supervision for generating structurally coherent and aesthetically plausible scenes with clean layouts. Following prior works [21, 29], we applied filters to remove layouts with uncommon object configurations or unnatural room proportions, ensuring a focus on typical and well-structured designs.

Procedurally Generated Scenes. To systematically enhance object diversity, particularly for small and decorative items, we generate the Inf3DLayout subset using the procedural generator Infinigen [24]. A key part of our curation involved carefully configuring the generator to produce plausible layouts for five common room types: bedrooms, living rooms, dining rooms, kitchens, and bathrooms. The generated houses were then programmatically partitioned into individual rooms. Finally, we applied a filtering step to remove rooms with abnormal layouts or spatial inconsistencies. This curated subset significantly increases the variety and granularity of object arrangements, covering numerous long-tail scenarios that are underrepresented in scan-based or design-based data.

The combination of these sources ensures that M3DLayout encompasses a wide spectrum of indoor environments, balancing realism, design integrity, and compositional diversity.

3.2. Structured Description Annotation

To support fine-grained text-conditioned layout generation, we annotate each 3D layout in the M3DLayout dataset with a comprehensive structured description. The annotation schema is designed to capture spatial and semantic information at multiple levels, comprising 3 key components.

Global Scene Description. This part captures the overall

Table 1. Quantitative analysis of three data sources (3D-FRONT, Matterport3D, Inf3DLayout) in M3DLayout.

Source	Scenes	Total Objects	Avg Objs/ Scene	Large Furniture	Small Objects	Small %
3D-FRONT	5,754	39,494	6.9	39,407	87	0.2%
Matterport3D	1,684	21,212	12.6	12,859	8,353	39.4%
Inf3DLayout	13,929	373,136	26.8	117,700	255,426	68.5%
	21,367	433,842	20.3	169,966	263,866	60.8%

Table 2. **Comparisons between existing 3D indoor scene datasets.** For the column of Layout Collection, RS denotes real scanned and PD denotes professionally designed. For the column of Variation in Object Sizes, “N/A” denotes “not available”, “L” and “S” denote Large and Small objects in the scene, respectively.

Dataset	Scenes	Objects	Layout Collection	Layout Complexity	Variation in Object Sizes	Structured Descriptions
SUN3D [33]	254	N/A	RS	Low	N/A	✗
SceneNN[15]	100	N/A	RS	Low	N/A	✗
Matterport3D[3]	1,684	N/A	RS	Medium	L-S	✗
ScanNet [6]	1,506	N/A	RS	Low	L-S	✗
Scan2CAD [1]	1,506	N/A	RS	Low	N/A	✗
OpenRooms [17]	1,068	97,607	RS	Low	N/A	✗
SceneNet [11]	57	3,699	PD	Low	N/A	✗
Structured3D [41]	N/A	N/A	PD	Low	N/A	✗
3D-FRONT [9]	5,754	N/A	PD	Low	L	✗
M3DLayout (Ours)	21,367	433,842	Mixture	High	L-S	✓

properties and organization of the scene. Each layout is labeled with the room type (e.g., dining area), stylistic attributes (e.g., modern style), and geometric features such as room shape and object density. It also includes high-level functional zoning (e.g., central dining zone with surrounding storage units) and global spatial patterns like symmetry (e.g., equal number of chairs on both sides of the table).

Large Furniture Description. We describe the presence and arrangement of major furniture pieces such as dining tables, shelves, consoles, and cabinets. The annotation includes both absolute positioning (e.g., “The shelving unit is placed opposite the entry door”) and relative spatial relations (e.g., “The console table is adjacent to the shelving unit”). A summary of large furniture composition is also provided for each room.

Small Object Description. This component focuses on decorative and functional small items like dishes, bowls, vases, books, and boxes. These are annotated based on their placement (e.g., on furniture surfaces or shelves) and distribution patterns (e.g., evenly distributed or symmetrically placed). Such fine-grained annotations support more detailed spatial reasoning and realistic scene generation.

The structured descriptions are generated through a multi-stage pipeline, as illustrated in Figure 2. As outlined in Section 3.1, we begin by generating the Inf3DLayout subset using a carefully configured procedural generation pipeline based on Infinigen. For layouts sourced from Matterport3D and the generated Inf3DLayout subset, we render top-down and side views, along with close-up images

highlighting the placement of small objects. These multi-view renders are subsequently processed by the GPT-4o model to produce the structured textual descriptions. For scenes from 3D-FRONT, we adopt a rule-based approach: after extracting object-level bounding boxes and semantic labels, we detect relative spatial relationships between objects and format this structured information into coherent descriptions using predefined templates. This methodology is suitable for 3D-FRONT due to the typically simpler and more regular spatial arrangements in its professionally designed scenes, allowing for accurate and comprehensive coverage via template-based generation. Finally, all automatically generated descriptions undergo a sampling-based manual review to ensure annotation quality.

3.3. Dataset Statistics and Analysis

M3DLayout encompasses a diverse collection of indoor environments, covering 26 scene categories. Among them, five core room types receive the most focus: bedroom, living room, dining room, kitchen, and bathroom. These constitute the most common residential spaces. In addition to these, the dataset includes functional areas such as office, entryway, closet, toilet, and balcony, as well as specialized spaces including gym, library, and home theater.

Data Composition and Source Characteristics. As detailed in Table 1, M3DLayout integrates 21,367 layouts with 433,842 object instances, averaging 20.3 objects per scene. The three data sources exhibit complementary characteristics: 3D-FRONT provides professionally designed layouts with clean structural regularity but limited small objects (0.2%); Matterport3D offers realistic scanned environments with moderate object density (12.6 objects/scene) and a balanced object distribution (39.4% small objects); while Inf3DLayout significantly enriches the dataset with high scene complexity (26.8 objects/scene) and abundant small objects (68.5%).

Object Distribution and Scene Complexity. Figure 3(a) shows the top 15 most frequent object categories, where small decorative items (e.g., Nature Shelf Trinkets, Book Stack) dominate the distribution, reflecting the dataset’s fine-grained annotation richness. The objects-per-scene distribution in Figure 3(b) reveals that M3DLayout covers a wide spectrum of scene complexities, from minimalist arrangements to densely populated environments. This variation enhances the generalization capability of trained models for real-world scenarios.

Comparative Advantages. As shown in Table 2, M3DLayout surpasses existing datasets in scale (21,367 scenes, 433k+ objects), layout complexity, and object size variation. Unlike datasets limited to either large furniture (L) or simple layouts, M3DLayout provides comprehensive coverage of both large and small objects (L-S) with structured textual descriptions, addressing a critical gap in cur-

Table 3. **Quantitative comparison.** Lower FID/KID ($\times 0.001$) and higher Clip Score indicate better synthesis quality. FID and KID are computed with respect to the real layouts from 3D-FRONT, Matterport, and Inf3DLayout. We train InstructScene following the public implementation. The optimal result is shown in bold, and the sub-optimal result is shown with an underline.

Method	FID ↓			KID ↓			CLIP-Score ↑
	3D-FRONT	Matterport	Inf3DLayout	3D-FRONT	Matterport	Inf3DLayout	
DiffuScene[29]	29.47	<u>98.03</u>	102.12	10.32	<u>47.92</u>	75.49	0.1982
InstructScene[18]	68.58	100.54	159.27	54.70	49.23	156.62	0.1944
Ours (DIFF-M3DLayout)	<u>57.64</u>	87.89	<u>70.85</u>	<u>36.80</u>	34.62	<u>50.94</u>	<u>0.2001</u>
Ours (AR-M3DLayout)	87.98	107.58	57.90	57.41	53.89	41.04	0.2026

rent data resources for detailed scene generation.

The broad coverage of room types, the detailed structured descriptions, and the variation in scene complexity make this dataset a valuable resource for downstream tasks such as layout prediction, text-conditioned scene synthesis, and embodied AI simulation.

4. Indoor Scene Generation

4.1. Problem Formulation

We formulate text-conditioned 3D indoor layout generation as a conditional generative problem. Given a scene-level natural language description c^{text} , the goal is to generate a structured 3D layout $x = \{o_i\}_{i=1}^N$ consisting of N objects.

Each object o_i is defined as a 3D oriented bounding box:

$$o_i = (c_i, x_i, y_i, z_i, w_i, h_i, d_i, \theta_i), \quad (1)$$

where c_i denotes the semantic class, (x_i, y_i, z_i) is the 3D object center, (w_i, h_i, d_i) are dimensions, and θ_i is the yaw rotation. The objective is to learn a conditional distribution $p_\theta(x | c^{\text{text}})$ that captures plausible spatial and semantic relationships among objects given the textual context.

We explore two complementary approaches for modeling this distribution: (1) a diffusion-based model, which learns to iteratively denoise noisy layout representations toward structured scenes, and (2) an autoregressive model, which generates scene objects sequentially in a structured order. Both methods are trained to reconstruct layouts consistent with the conditioning text while maintaining spatial coherence and physical plausibility.

4.2. Diffusion-based Model Architecture

Our diffusion model follows a denoising diffusion probabilistic framework similar to DiffuScene. The forward process gradually perturbs layout representations x_0 with Gaussian noise following a fixed variance schedule β_t , while the reverse process is parameterized by a neural network ϵ_θ that predicts the noise conditioned on the text input:

$$p_\theta(x_{t-1} | x_t, c^{\text{text}}) = \mathcal{N}(x_{t-1}; \mu_\theta(x_t, t, c^{\text{text}}), \Sigma_\theta(x_t, t)). \quad (2)$$

The model is trained with a noise-prediction objective

$$\mathcal{L}_{\text{DM}} = \mathbb{E}_{x_0, c^{\text{text}}, t, \epsilon} [\|\epsilon - \epsilon_\theta(x_t, t, c^{\text{text}})\|_2^2], \quad (3)$$

and extra regularization terms such as an IoU-based collision penalty to encourage physically valid object arrangements. More details are in the supplementary material.

4.3. Autoregressive Model Architecture

To complement the diffusion paradigm, we further employ an autoregressive transformer model that generates scene layouts sequentially. Given the textual condition c^{text} , the model factorizes the layout likelihood as

$$p_\theta(x | c^{\text{text}}) = \prod_{i=1}^N p_\theta(o_i | o_{<i}, c^{\text{text}}), \quad (4)$$

where each object o_i is predicted conditioned on the previously generated objects and the text.

The autoregressive model is implemented with a transformer encoder architecture. Text tokens and previously generated object embeddings form a unified input sequence for context-aware sequential prediction. The training objective minimizes the negative log-likelihood of object parameters under the predicted conditional distribution. Further details are provided in the supplementary material.

5. Experiments

5.1. Experimental Settings

We evaluate the effectiveness of autoregressive and diffusion-based methods on the proposed dataset, assessing the plausibility and controllability of generated 3D layouts. Detailed train/val splits and implementation details are in the supplementary material.

Baselines. We compare our method with two state-of-the-art scene generation approaches: DiffuScene [29] and InstructScene [18], both of which are text-driven methods based on diffusion models. More details are provided in the supplementary material.

Metrics. Following prior works [18, 29], we adopt Fréchet Inception Distance (FID) [13] and Kernel Inception Distance (KID) [2], and the CLIP-Score [12] to measure the

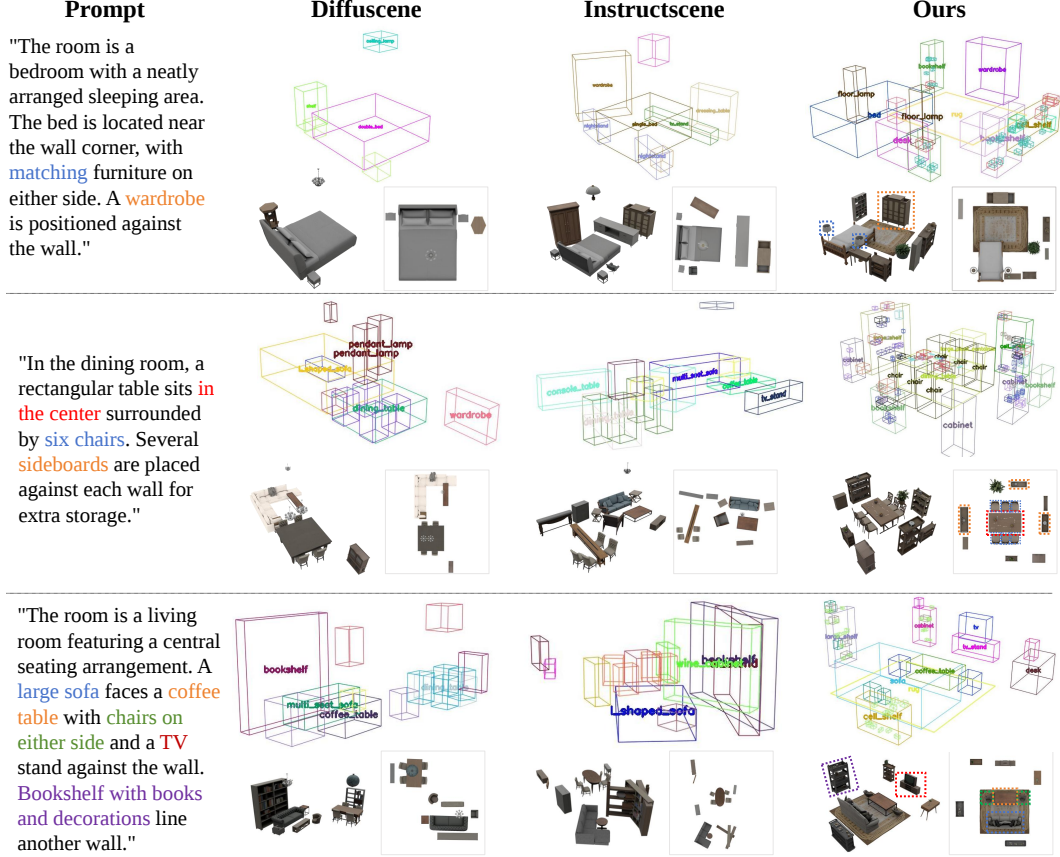


Figure 4. **Qualitative comparison of different methods on diverse room types.** From top to bottom: bedroom, dining room, and living room generation results. Each row shows the input prompt and generated layouts from Diffuscene, Instructscene, and our method. Trained on the M3DLayout dataset, our method produces richer layout details from text descriptions.

controllability and fidelity, respectively. We employ a self-designed object retrieval method to realize instance filling from layout to scene and rendering. More details are provided in the supplementary material.

5.2. Main Results

Quantitative Comparison.

The quantitative comparison results are shown in Table 3, where “DIFF-M3DLayout” and “AR-M3DLayout” represent diffusion- and autoregressive-based methods, respectively. In Table 3, the horizontal axis represents ground-truth renderings (real images) from three different datasets. The same set of 1,500 prompts is used as conditioning input for all comparative methods to generate layouts, which are then rendered as synthetic images and used together with real images to compute FID and KID. More details can be found in the supplementary material.

As demonstrated, for FID and KID, the proposed DIFF-M3DLayout drastically outperforms the state-of-the-art, achieving improvements of 10%–32% on the reference Matterport dataset and Inf3DLayout datasets, and

AR-M3DLayout achieves 34%–44% improvements on Inf3DLayout dataset, which demonstrates superior generalization compared with DiffuScene and InstructScene. On 3D-FRONT, however, it falls behind DiffuScene on these metrics. This is primarily because the number of objects per scene in 3D-FRONT typically ranges from 5 to 12, whereas our method generates scenes with more than 12 objects in most cases, causing a mismatch in scene complexity distribution. As a result, the visual statistics of our generated layouts deviate more from the ground truth (3D-FRONT), which negatively impacts FID and KID. Conversely, the richer object counts and variety in the scenes generated by our model provide another advantage of our method relative to the multi-source datasets M3DLayout, which is also confirmed by the visualizations in Figure 4. Moreover, our method surpasses the baselines in CLIP-Score, demonstrating enhanced controllability and a stronger alignment between the generated scenes and the given prompts.

Qualitative Comparison. The proposed DIFF-M3DLayout is employed for qualitative comparison with diverse methods, with the results in the bedroom, din-

Shared Description: In this dining room, a farmhouse-style table takes center stage. Chairs are arranged around it with a side table placed near the entrance.

Simple: Do not include small decorative items.

Basic: (No extra description)

Detailed: Add rich details to the objects.

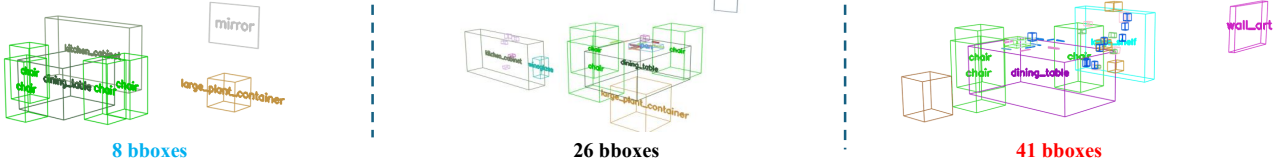


Figure 5. **Density controllability in layout generation with different input texts.** The first row presents input prompts for our layout generation model, showcasing variations in objects density from low to high, with minor changes in the last sentence. The second row illustrates the corresponding output results generated by our model, which adapt based on the prompt density.

ing room, and living room shown in Figure 4. Across all settings, the proposed DIFF-M3DLayout demonstrates improved semantic controllability and visual fidelity compared to DiffuScene and InstructScene. Specifically, DiffuScene can generate scenes that appear visually neat at first glance but struggles to produce small objects and exhibits limited prompt controllability, as illustrated by the living room case. InstructScene, on the other hand, fails to accurately model spatial relationships among instances, leading to disorderly object placement, as shown in the dining room case. By contrast, the proposed DIFF-M3DLayout generates precise and diverse objects in accordance with the prompt semantics—such as the small items on the shelf and six chairs in the dining room case—while effectively capturing 3D spatial arrangements, exemplified by “The bed is located near the wall corner” in the bedroom case.

We perform the ablation experiments using the diffusion-based method to validate the effectiveness of a single training dataset. More details and results are provided in the supplementary material.

5.3. Density Controllability

We verify the controllability of object density in our layout generation model using different input prompts. As shown in Figure 5, the first row displays the different types of input prompts, which only differ in the final sentence. These prompts adjust the density of objects within the layout, from a minimal setup to one with richer details. This highlights the model’s ability to control scene density based on the granularity of the input description, thus offering flexibility in layout customization for various applications.

5.4. User Case Study

We conducted a perceptual study to evaluate the quality of our generated layouts against DiffuScene and InstructScene. We recruited 42 participants to rate 15 scenes across three room types (dining room, bedroom, and living room). For each scene, participants were shown a text description, a top-down rendering, and the generated layout for each of the three methods, rating them on a 1-to-5 scale

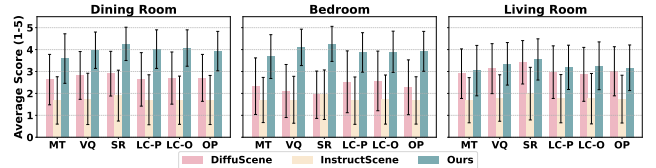


Figure 6. **User case study results.** The charts compare our method against DiffuScene and InstructScene across diverse rooms. Bars represent the average user score for six metrics: Match with Text (MT), Visual Quality (VQ), Scene Richness (SR), Layout Coherence (Position) (LC-P), Layout Coherence (Orientation) (LC-O), and Overall Preference (OP).

across six metrics. All participants were volunteers without compensation. The overall results are summarized in Figure 6. As shown, our method outperformed both baselines in the vast majority of metrics and room categories. The most significant advantage was observed in the Scene Richness metric. The performance gap was less pronounced in living rooms, likely because key layouts are harder to assess from the top-down view used compared to the iconic beds or dining sets in other rooms. These results confirm that users perceive our generated layouts as more detailed, coherent, and of higher quality.

6. Conclusion

We introduced M3DLayout, a large-scale, multi-source dataset for 3D indoor layout generation from structured text descriptions. It integrates real-world scans, professional designs, and procedurally generated scenes. Each layout is paired with structured descriptions that cover global scene summaries, relational placements of large furniture, and fine-grained arrangements of small objects. This multi-source, richly annotated structure enables the learning of diverse spatial and semantic patterns across a wide variety of indoor environments. We established a benchmark using both a text-conditioned diffusion model and a text-conditioned autoregressive model, and experimental results validate that training on our dataset enhances the diversity and detail of generated layouts, with the Inf3DLayout sub-

set in particular enabling more complex and richly annotated scenes. Despite these advances, our dataset has certain limitations. The structured descriptions, while comprehensive, are partly generated via language models and may contain noise. We hope M3DLayout serves as a valuable resource for advancing research in text-driven 3D scene synthesis.

M3DLayout: A Multi-Source Dataset of 3D Indoor Layouts and Structured Descriptions for 3D Generation

Supplementary Material

7. Supplementary Material Overview

This supplementary document provides additional technical details and visualization results to support the main paper. The supplementary is organized into four sections:

Sec. 8 Implementation Details. We provide detailed explanations of the dataset splitting strategy in 8.1, both diffusion and autoregressive based model details in 8.2, training details in 8.3, compared baselines details in 8.4 and metric details in 8.5.

Sec. 9 More Visualization Results. We provide additional examples that further demonstrate the quality of our dataset, the fidelity of our generative model, and its coherence with textual inputs.

Sec. 10 Ablation. We conduct some ablation studies to show that models trained on a single dataset overfit to its distribution and generalize poorly, whereas training on the multi-source M3DLayout dataset yields consistently balanced, realistic, and controllable layouts across diverse data types.

Sec. 11 Object Retrieval. We present a scalable and self-developed layout-to-scene object retrieval pipeline that accurately maps generated layouts to 3D assets, facilitating reliable metric evaluation and high-fidelity visualizations.

8. Implementation Details

8.1. Dataset Split

Our experiments are conducted on a subset of the M3DLayout dataset containing 15,080 scenes. This subset is randomly divided into 12,062 layouts for training and 3,018 layouts for validation. For the ablation studies in Table 5, models are additionally trained on three independent datasets with the following splits (training/validation): 4,603/1,151 for 3D-FRONT, 1,347/337 for Matterport3D, and 6,112/1,530 for Inf3DLayout. For all evaluations, and to ensure a fair comparison with prior work, we do *not* use the scene description texts provided in M3DLayout. Instead, we generate 500 scene descriptions each for *bedroom*, *dining room*, and *living room* using GPT-4o, resulting in a total of 1,500 test descriptions. This unified test set is used across all experiments for fairness.

8.2. Model Details

Diffusion-based Model. Our diffusion-based model represents a scene as a sequence of N objects, where the maximum number of objects N is fixed at 120. Padding objects with a PAD category label is applied to maintain a consistent sequence length. The sequence is ordered by the (x, z, y) coordinates of the bounding box centers, following recent best practices in the 3D part-level generation community for layout generation. Each object is parameterized by a concatenated vector $[l_i, s_i, \cos \theta_i, \sin \theta_i, c_i]$, which includes its location ($l_i \in \mathbb{R}^3$), size ($s_i \in \mathbb{R}^3$), orientation ($\theta_i \in \mathbb{R}$, encoded using $\cos \theta_i$ and $\sin \theta_i$), and category label ($c_i \in \mathbb{R}^C$).

The denoiser is a UNet-based architecture built upon 1D convolutions with skip connections. It independently encodes and predicts the distinct attributes of the object. We use a pre-trained BERT encoder to extract text embeddings z from the input description. This language guidance is injected into the denoiser network through cross-attention layers, enabling the network to predict the noise ϵ_ϕ conditioned on the timestep t and the text embedding z .

The training objective consists of the scene loss L_{sce} , which minimizes the difference between the actual and predicted noise, supplemented with an IoU regularization term L_{iou} to penalize object intersections and encourage physically plausible arrangements.

Autoregressive Model. Our autoregressive model shares the same scene representation as our diffusion-based model for consistency. The autoregressive approach does not have a strict upper limit on the sequence length, but we still utilize data containing up to 120 objects per scene for training. A scene is represented as a sequence of objects, bookended by special start-of-sequence (SOS) and end-of-sequence (EOS) tokens.

The input text description is similarly encoded into text embeddings using a pre-trained BERT model. At each generation step, the input sequence to the model consists of the text token embeddings and the embeddings of all previously generated objects. A Transformer encoder serves as the backbone network. The attributes of each previously generated object are first projected by an MLP into an embedding space. The network then attends to this combined sequence to predict the parameters of the next object.

The training objective is the scene loss L_{sce} , which minimizes the negative log-likelihood of the ground-truth object parameters given the input context.

8.3. Training Details.

Both our diffusion and autoregressive based models are trained for 30k epochs using the Adam optimizer. For the diffusion model, we use a learning rate of 2×10^{-4} with a step-wise decay factor of 0.5 every 10k steps and a linear noise schedule. For the autoregressive model, we use a learning rate of 1×10^{-4} .

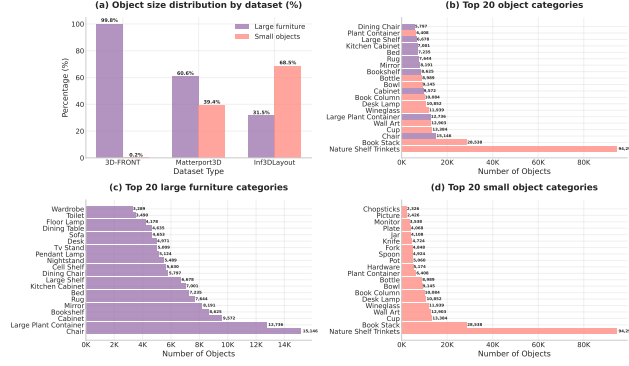


Figure 7. **Object distribution statistics of the M3DLayout dataset.** (a) Size distribution (large/small) of objects by source. (b) Overall ranking of the top 20 object categories. (c-d) Rankings for large and small objects, respectively.

8.4. Compared Baselines

We compare our method with two state-of-the-art scene generation approaches: (1) DiffuScene [29], a diffusion model for 3D indoor scene synthesis that denoises unordered object attributes to produce physically plausible layouts. (2) InstructScene [18], a graph diffusion model that integrates a semantic graph with a layout decoder to synthesize 3D indoor scenes from natural language instructions. Both methods allow for the conditioning on text prompts. For inference with DiffuScene, we employ the officially released model weights for the bedroom, dining room, and living room, and follow the public implementation to train InstructScene on the same room types.

8.5. Metrics

Following prior works [18, 29], we adopt Fréchet Inception Distance (FID) [13] and Kernel Inception Distance (KID) [2] to quantify the fidelity of scenes synthesized from layouts by measuring the similarity between generated and ground-truth top-down renderings. Meanwhile, we employ the CLIP-Score [12] to evaluate the controllability of generated layouts by computing the cosine similarity between CLIP-encoded features of the generated renderings and the given prompts. To this end, we first employ a text2mesh model [32] to generate the required object instances and then retrieve both the ground-truth and synthesized scenes conditioned on the layout.

As for the calculation of FID and KID, the real images are typically taken from the training set; however, since our method and the baselines are trained on different datasets with varying data distributions, direct comparison would be unfair. To address this, we select different datasets as the source of real images, allowing for both a fair comparison and an evaluation of fidelity and controllability.

9. More Visualization Results

9.1. Dataset Layout Visualization

We provide diverse types of 3D scene data from our M3DLayout dataset in Figure 8, which includes scenes from CAD designs sourced from the 3D-FRONT dataset at the first row, scenes derived from real-world scans, specifically from the Matterport3D dataset at the second row and procedurally generated scenes Inf3DLayout from Infinigen at the third row. These images demonstrate the flexibility of the M3DLayout dataset in representing a wide spectrum of interior environments, from synthetic CAD designs to real-world captures and generative models.

9.2. Generated Layout Visualization

We visualize more generated layouts by both our diffusion and autoregressive model trained on the M3DLayout dataset, involving bedroom, living room, and dining room in Figure 9 and Figure 10 respectively. From the table, it is evident that our method achieves remarkable performance in both layout coherence and the richness of objects in the generated scenes. We also provide more visualization to indicate strong text coherence of our generated layouts in Figure 11. These qualitative visualizations further highlight that our method surpasses prior state-of-the-art approaches in both fidelity and controllability.

10. Ablation Study

We perform the ablation experiments using the diffusion- and autoregressive-based methods to validate the effectiveness of a single training dataset and report the results in Table 5. As shown in the table, for the diffusion-based method, when the training data and ground truth come from the same dataset, the model trained on a single dataset achieves the best FID and KID compared with the other two models. However, its performance drops significantly when evaluated on data from different datasets. This indicates that, while models trained on a single dataset can effectively fit the distribution of that dataset, they struggle to generalize to varied data. For example, provided textual guidance, a model trained on the professional CAD designs dataset (3D-FRONT) encounters difficulties in generating scenes that align with real-world scans (Matterport) or procedurally generation (Inf3DLayout) dataset. Furthermore, the

Table 4. Comparison of generation efficiency with state-of-the-art methods.

Method	Generation Time per Scene (s)	Generation Time per BBox (s)	Average BBoxes per Scene
Diffuscene	17.4067	1.7674	9.849
InstructScene	<u>1.3438</u>	<u>0.1280</u>	10.502
Ours (DIFF-M3DLayout)	30.18	1.639	18.4
Ours (AR-M3DLayout)	0.2348	0.01587	<u>14.8</u>

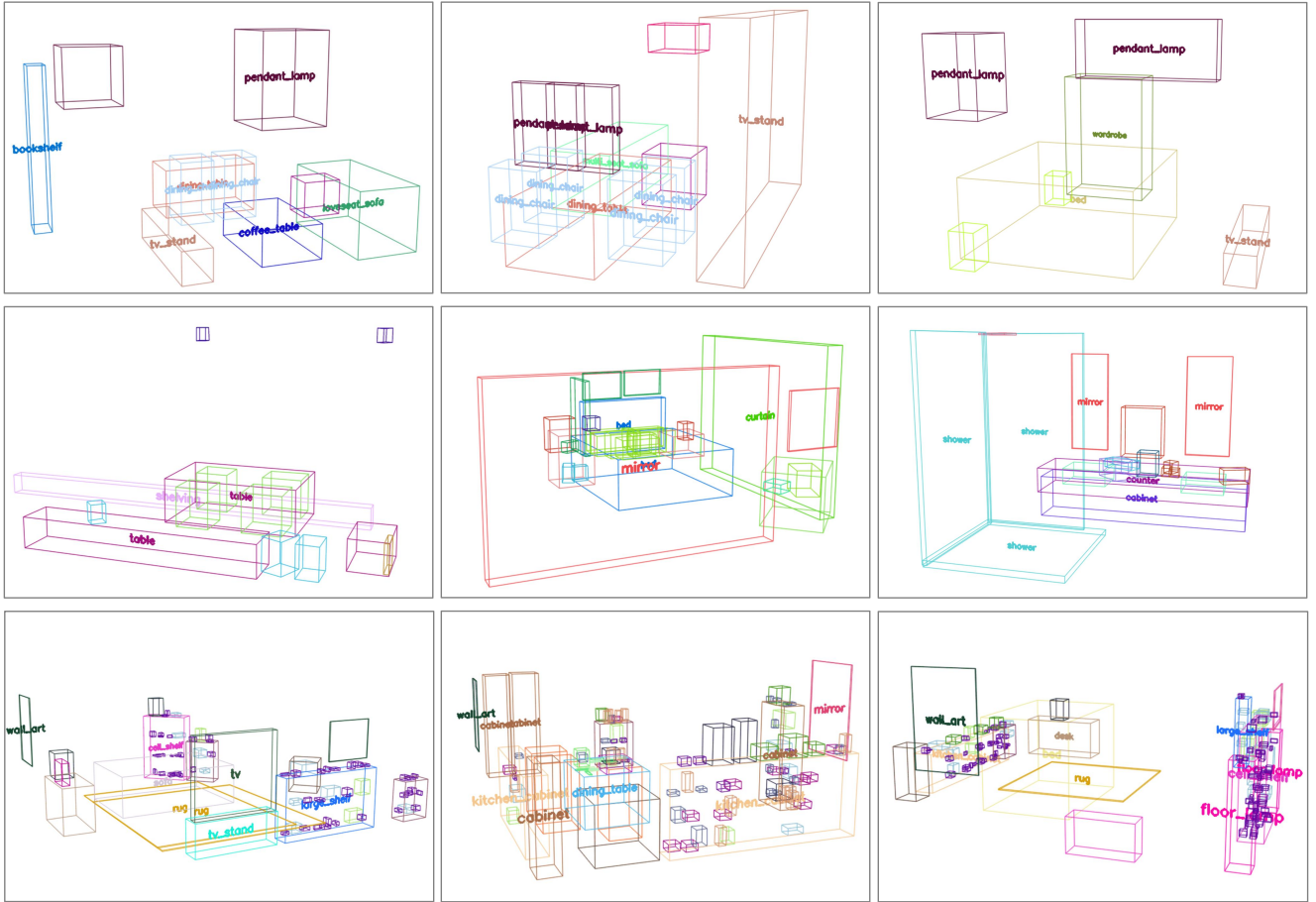


Figure 8. Samples from the M3DLayout dataset. The first row shows scenes from CAD designs (3D-FRONT), the second row from real-world scans (Matterport3D), and the third row from procedurally generated scenes (Infinigen).

autoregressive-based method exhibits the same characteristics. In contrast, these methods trained on the multi-source M3DLayout dataset achieves balanced performance across data types, producing more realistic and controllable layouts (see Table 3 in Section 5.2 of the present work for comprehensive data).

11. Object Retrieval

To effectively visualize the generated layouts and meet the evaluation requirements, such as FID (Fréchet Incep-

tion Distance), KID (Kernel Inception Distance), and CLIP score, we present a simple, effective and scalable pipeline for layout-to-scene object retrieval.

In Figure 12, we first build our retrieval dataset by constructing huge amounts of *prompts* for 95 object categories (See details in Table 6), which covers all objects for our dataset, as input for *Text-to-3D generation model* (TRELIS [32]). By delicately designing our prompts, we can obtain 3D assets with different scales, textures and application scenarios, which can be generalizable to handle with intri-

Table 5. **Ablation studies of diffusion-based methods trained on different datasets.** Lower FID/KID ($\times 0.001$) and higher Clip Score indicate better synthesis quality. FID and KID are computed with respect to the real layouts from 3D-FRONT, Matterport, and Inf3DLayout.

Method		FID ↓			KID ↓			CLIP-Score ↑
		3D-FRONT	Matterport	Inf3DLayout	3D-FRONT	Matterport	Inf3DLayout	
Diffusion-based	Ours (3D-FRONT)	27.33	83.88	110.98	10.59	21.80	83.45	0.2083
	Ours (Matterport)	81.31	69.61	114.58	46.82	18.41	94.45	0.1916
	Ours (Inf3DLayout)	93.51	115.07	54.36	55.67	55.53	34.95	0.1969
Autoregressive-based	Ours (3D-FRONT)	30.02	85.08	117.39	12.09	24.46	95.11	0.2056
	Ours (Matterport)	73.36	61.08	132.54	45.11	12.58	122.21	0.1959
	Ours (Inf3DLayout)	105.24	122.19	60.73	70.81	65.21	43.21	0.2026

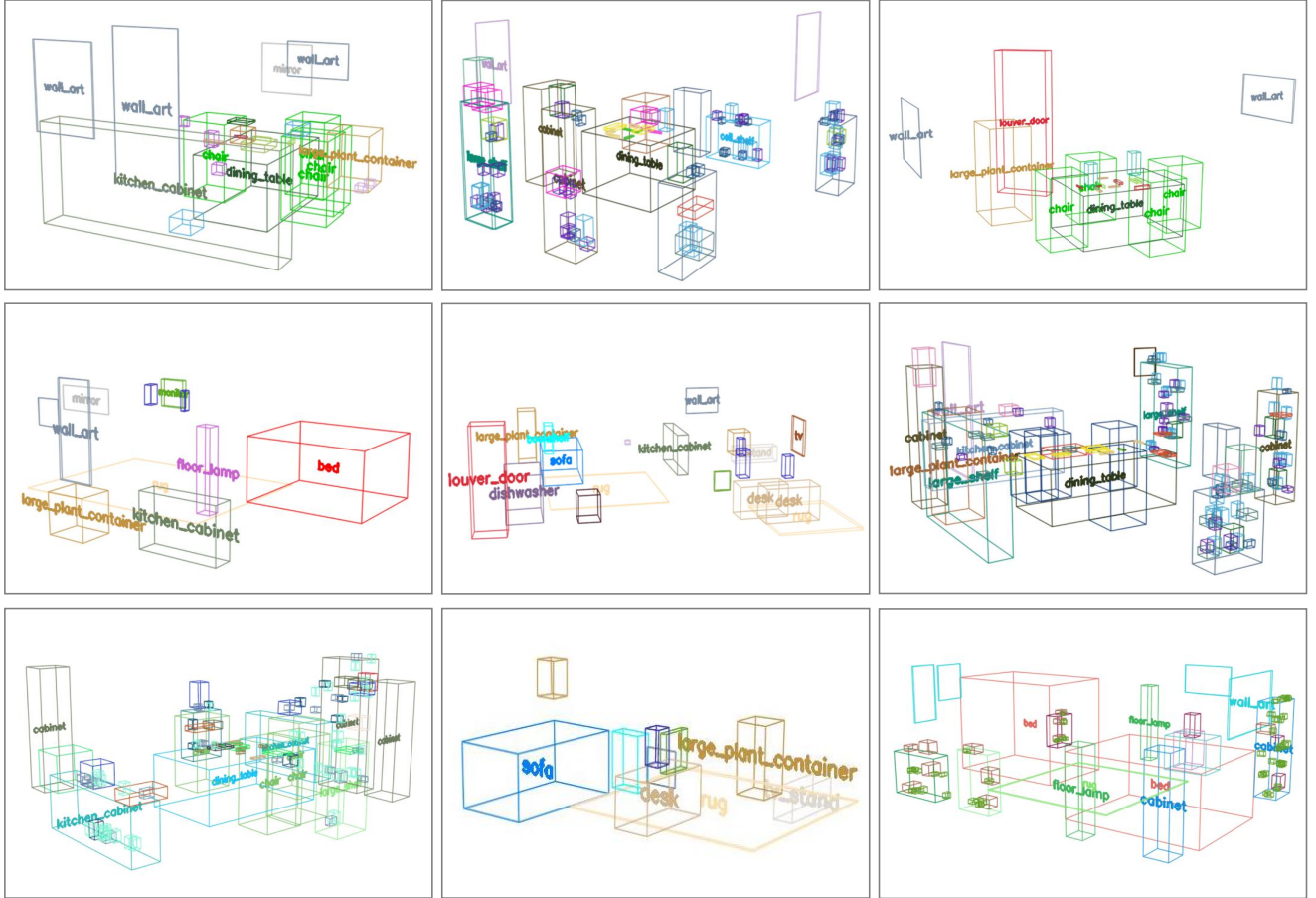


Figure 9. Generated layouts by our model trained on the M3DLayout dataset, visualized for randomly selected bedroom, living room, and dining room.

cate object retrieval process.

In the *Object Selection* phase, the retrieved object, such as a bed, undergoes attribute extraction through the *Attribute Solver*. This solver precomputes attributes like width, height, and depth ratios for each object in the retrieval dataset, and extracts scalar and categorical information from bounding box (BBox) of each object in the generated layout. The BBox, along with additional dataset information, is passed to a *Shape & Category Similarity Solver*

to match the most appropriate object.

Finally, after iterating all objects in the scene, all best matching objects are chosen and their properties (such as translations, sizes, and rotation) are determined, culminating in the retrieval of the desired 3D scene. This multi-step process ensures accurate retrieval and selection of 3D objects for following applications.

After retrieving successfully, the visualizations provided in this Figure 13 aim to assess both the quality of the gen-

Category	Objects (95 in total)
Lighting	lighting, ceiling_lamp, pendant_lamp, floor_lamp, desk_lamp, fan
Tables	table, coffee_table, console_table, corner_side_table, round_end_table, dining_table, dressing_table, side_table, nightstand, desk, tv_stand
Seating	seating, chair, armchair, lounge_chair, chinese_chair, dining_chair, dressing_chair, stool, sofa, loveseat_sofa, l_shaped_sofa, multi_seat_sofa
Beds	bed, kids_bed
Shelves & Book storage	shelf, shelving, large_shelf, cell_shelf, bookshelf, book, book_column, book_stack, nature_shelf_trinkets
Cabinets & Wardrobes	cabinet, kitchen_cabinet, children_cabinet, wardrobe, wine_cabinet
Appliances & Electronics	appliances, microwave, oven, beverage_fridge, tv, monitor, tv_monitor
Kitchen & Tableware	pan, pot, plate, bowl, cup, bottle, can, jar, wineglass, chopsticks, knife, fork, spoon, food_bag, food_box, fruit_container
Bathroom fixtures	bathtub, shower, sink, standing_sink, toilet, toilet_paper, toiletry, faucet, towel
Doors, Windows & Coverings	glass_panel_door, lite_door, window, blinds, curtain, vent
Hardware & Controls	hardware, handle, light_switch
Decor	plant, large_plant_container, plant_container, vase, wall_art, picture, mirror, statue, basket, balloon, cushion, rug, decoration
Containers & Waste	bag, box, container, clutter, trashcan
Architecture & Elements	counter, fireplace, pipe, furniture
Clothes	clothes
Spaces	kitchen_space
Gym & Misc	gym_equipment

Table 6. **Category list of retrieval objects.** Our retrieval dataset includes 95 objects which covers nearly all common indoor objects.

- William Yang Wang. LayoutGPT: Compositional Visual Planning and Generation with Large Language Models. *Advances in Neural Information Processing Systems*, 36: 18225–18250, 2023. 2
- [9] Huan Fu, Bowen Cai, Lin Gao, Ling-Xiao Zhang, Jiaming Wang, Cao Li, Qixun Zeng, Chengyue Sun, Rongfei Jia, Binqiang Zhao, and Hao Zhang. 3D-FRONT: 3D Furnished Rooms with layOuts and semaNTics. In *2021 IEEE/CVF International Conference on Computer Vision (ICCV)*, pages 10913–10922, 2021. 2, 3, 4, 5
- [10] Lin Gao, Jia-Mu Sun, Kaichun Mo, Yu-Kun Lai, Leonidas J. Guibas, and Jie Yang. SceneHGN: Hierarchical Graph Networks for 3D Indoor Scene Generation with Fine-Grained Geometry, 2023. 3
- [11] Ankur Handa, Viorica Patraucean, Vijay Badrinarayanan, Simon Stent, and Roberto Cipolla. Understanding real world indoor scenes with synthetic data. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4077–4085, 2016. 5
- [12] Jack Hessel, Ari Holtzman, Maxwell Forbes, Ronan Le Bras, and Yejin Choi. Clipscore: A reference-free evaluation metric for image captioning. *arXiv preprint arXiv:2104.08718*, 2021. 6, 2
- [13] Martin Heusel, Hubert Ramsauer, Thomas Unterthiner, Bernhard Nessler, and Sepp Hochreiter. Gans trained by a two time-scale update rule converge to a local nash equilibrium. *Advances in neural information processing systems*, 30, 2017. 6, 2
- [14] Lukas Höllein, Ang Cao, Andrew Owens, Justin Johnson, and Matthias Nießner. Text2Room: Extracting Textured 3D Meshes from 2D Text-to-Image Models, 2023. 2
- [15] Binh-Son Hua, Quang-Hieu Pham, Duc Thanh Nguyen, Minh-Khoi Tran, Lap-Fai Yu, and Sai-Kit Yeung. SceneNN: A Scene Meshes Dataset with aNnotations. In *2016 Fourth International Conference on 3D Vision (3DV)*, pages 92–101, Stanford, CA, USA, 2016. IEEE. 3, 5
- [16] Zehuan Huang, Yuan-Chen Guo, Xingqiao An, Yunhan Yang, Yangguang Li, Zi-Xin Zou, Ding Liang, Xihui Liu, Yan-Pei Cao, and Lu Sheng. MIDI: Multi-Instance Diffusion for Single Image to 3D Scene Generation, 2024. 2
- [17] Zhengqin Li, Ting-Wei Yu, Shen Sang, Sarah Wang, Meng Song, Yuhua Liu, Yu-Ying Yeh, Rui Zhu, Nitesh Gundavarapu, Jia Shi, Sai Bi, Zexiang Xu, Hong-Xing Yu, Kalyan Sunkavalli, Miloš Hašan, Ravi Ramamoorthi, and Manmohan Chandraker. OpenRooms: An End-to-End Open Framework for Photorealistic Indoor Scene Datasets, 2021. 3, 5
- [18] Chenguo Lin and Yadong Mu. InstructScene: Instruction-Driven 3D Indoor Scene Synthesis with Semantic Graph Prior, 2024. 2, 6
- [19] Gabrielle Littlefair, Niladri Shekhar Dutt, and Niloy J. Mitra. FlairGPT: Repurposing LLMs for Interior De-

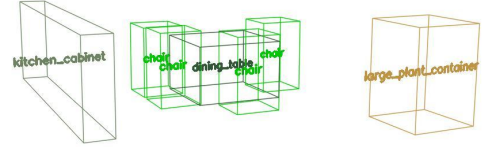
- signs. *Computer Graphics Forum*, 44(2):e70036, 2025. arXiv:2501.04648 [cs]. 4
- [20] Hao Ouyang, Kathryn Heal, Stephen Lombardi, and Tiancheng Sun. Text2Immersion: Generative Immersive Scene with 3D Gaussians, 2023. 2
- [21] Despoina Paschalidou, Amlan Kar, Maria Shugrina, Karsten Kreis, Andreas Geiger, and Sanja Fidler. ATISS: Autoregressive Transformers for Indoor Scene Synthesis. In *Advances in Neural Information Processing Systems*, pages 12013–12026. Curran Associates, Inc., 2021. 2, 3, 4
- [22] Pulak Purkait, Christopher Zach, and Ian Reid. SG-VAE: Scene Grammar Variational Autoencoder to Generate New Indoor Scenes. In *Computer Vision – ECCV 2020*, pages 155–171. Springer International Publishing, Cham, 2020. 3
- [23] Alexander Raistrick, Lahav Lipson, Zeyu Ma, Lingjie Mei, Mingzhe Wang, Yiming Zuo, Karhan Kayan, Hongyu Wen, Beining Han, Yihan Wang, Alejandro Newell, Hei Law, Ankit Goyal, Kaiyu Yang, and Jia Deng. Infinite Photorealistic Worlds Using Procedural Generation. In *2023 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 12630–12641, 2023. 3
- [24] Alexander Raistrick, Lingjie Mei, Karhan Kayan, David Yan, Yiming Zuo, Beining Han, Hongyu Wen, Meenal Parakh, Stamatis Alexandropoulos, Lahav Lipson, Zeyu Ma, and Jia Deng. Infinigen Indoors: Photorealistic Indoor Scenes using Procedural Generation. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 21783–21794, 2024. 3, 4
- [25] Jonas Schult, Sam Tsai, Lukas Höllein, Bichen Wu, Jialiang Wang, Chih-Yao Ma, Kunpeng Li, Xiaofang Wang, Felix Wimbauer, Zijian He, Peizhao Zhang, Bastian Leibe, Peter Vajda, and Ji Hou. ControlRoom3D: Room Generation Using Semantic Proxy Rooms. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 6201–6210, 2024. 2
- [26] Shaocong, Lihe Dong, Xiao Ding, Yaokun Chen, Yuxin Li, Yucheng WANG, Qi Wang, Jaehyeok WANG, Chenjian Kim, Zhanpeng Gao, Zibin Huang, Tianfan Wang, Dan Xue, and Xu. From one to more: Contextual part latents for 3d generation. 2025. 2
- [27] Shuran Song, Fisher Yu, Andy Zeng, Angel X. Chang, Manolis Savva, and Thomas Funkhouser. Semantic Scene Completion from a Single Depth Image. In *2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 190–198, Honolulu, HI, 2017. IEEE. 3
- [28] Xiaohao Sun, Divyam Goel, and Angel X. Chang. SemLayoutDiff: Semantic Layout Generation with Diffusion Model for Indoor Scene Synthesis, 2025. arXiv:2508.18597 [cs] version: 1. 3
- [29] Jiapeng Tang, Yinyu Nie, Lev Markhasin, Angela Dai, Justus Thies, and Matthias Nießner. DiffuScene: Denoising Diffusion Models for Generative Indoor Scene Synthesis. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 20507–20518, 2024. 2, 3, 4, 6
- [30] Xinpeng Wang, Chandan Yeshwanth, and Matthias Nießner. SceneFormer: Indoor Scene Generation with Transformers. In *2021 International Conference on 3D Vision (3DV)*, pages 106–115, 2021. 2, 3
- [31] Yian Wang, Xiaowen Qiu, Jiageng Liu, Zhehuan Chen, Jiting Cai, Yufei Wang, Tsun-Hsuan Wang, Zhou Xian, and Chuang Gan. Architect: Generating Vivid and Interactive 3D Scenes with Hierarchical 2D Inpainting. *Advances in Neural Information Processing Systems*, 37:67575–67603, 2024. 4
- [32] Jianfeng Xiang, Zelong Lv, Sicheng Xu, Yu Deng, Ruicheng Wang, Bowen Zhang, Dong Chen, Xin Tong, and Jiaolong Yang. Structured 3d latents for scalable and versatile 3d generation. In *Proceedings of the Computer Vision and Pattern Recognition Conference*, pages 21469–21480, 2025. 2, 3
- [33] Jianxiong Xiao, Andrew Owens, and Antonio Torralba. Sun3d: A database of big spaces reconstructed using sfm and object labels. In *Proceedings of the IEEE international conference on computer vision*, pages 1625–1632, 2013. 5
- [34] Xinhao Yan, Jiachen Xu, Yang Li, Changfeng Ma, Yunhan Yang, Chunshi Wang, Zibo Zhao, Zeqiang Lai, Yunfei Zhao, Zhuo Chen, et al. X-part: high fidelity and structure coherent shape decomposition. *arXiv preprint arXiv:2509.08643*, 2025. 2
- [35] Bangbang Yang, Wenqi Dong, Lin Ma, Wenbo Hu, Xiao Liu, Zhaopeng Cui, and Yuewen Ma. DreamSpace: Dreaming Your Room Space with Text-Driven Panoramic Texture Propagation. In *2024 IEEE Conference Virtual Reality and 3D User Interfaces (VR)*, pages 650–660. IEEE Computer Society, 2024. 2
- [36] Yue Yang, Fan-Yun Sun, Luca Weihs, Eli Vanderbilt, Alvaro Herrasti, Winsun Han, Jiajun Wu, Nick Haber, Ranjay Krishna, Lingjie Liu, Chris Callison-Burch, Mark Yatskar, Aniruddha Kembhavi, and Christopher Clark. Holodeck: Language Guided Generation of 3D Embodied AI Environments. In *2024 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 16277–16287, 2024. 2, 4
- [37] Yunhan Yang, Yufan Zhou, Yuan-Chen Guo, Zi-Xin Zou, Yukun Huang, Ying-Tian Liu, Hao Xu, Ding Liang, Yan-Pei Cao, and Xihui Liu. Omnipart: Part-aware 3d generation with semantic decoupling and structural cohesion. *arXiv preprint arXiv:2507.06165*, 2025. 2
- [38] Guangyao Zhai, Evin Pinar Örneç, Shun-Cheng Wu, Yan Di, Federico Tombari, Nassir Navab, and Benjamin Busam. CommonScenes: Generating Commonsense 3D Indoor Scenes with Scene Graph Diffusion. *Advances in Neural Information Processing Systems*, 36:30026–30038, 2023. 2, 3
- [39] Guangyao Zhai, Evin Pinar Örneç, Dave Zhenyu Chen, Ruotong Liao, Yan Di, Nassir Navab, Federico Tombari, and Benjamin Busam. EchoScene: Indoor Scene Generation via Information Echo Over Scene Graph Diffusion. In *Computer Vision – ECCV 2024*, pages 167–184, Cham, 2025. Springer Nature Switzerland. 2, 3
- [40] Genghao Zhang, Yuxi Wang, Chuanchen Luo, Shibiao Xu, Junran Peng, Zhaoxiang Zhang, and Man Zhang. FurniScene: A Large-scale 3D Room Dataset with Intricate Furnishing Scenes, 2024. 3
- [41] Jia Zheng, Junfei Zhang, Jing Li, Rui Tang, Shenghua Gao,

and Zihan Zhou. Structured3D: A Large Photo-realistic Dataset for Structured 3D Modeling, 2020. [2](#), [3](#), [5](#)

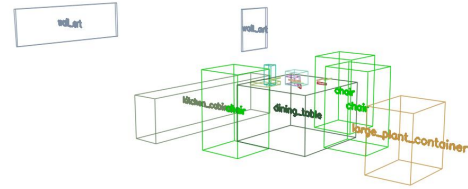
Input

A cozy dining room with a square table **in the center**. Chairs rest on **all sides**, and a **planter** adds greenery to one corner.

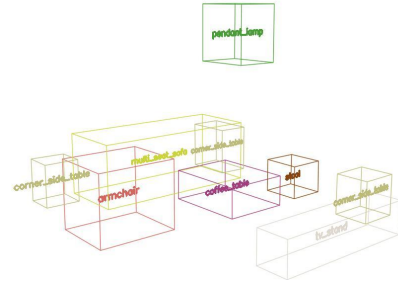
Output



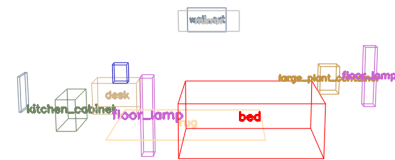
A dining room with a circular table has **evenly placed chairs** around it. A **cabinet** lines the wall, enhancing the dining experience.



This living room has a cozy corner with a sectional **sofa** and a small round **coffee table**. A **tv stand** lines the wall opposite the sofa, and a **side table** is placed beside it.



In this bedroom, the **bed** is against the back wall, with a **floor lamp** at its foot. A **small desk and table lamp** create a work area.



The room is a bedroom and includes a **single bed** pushed into a corner. Across the room, a **cabinet** provide a place to store things.

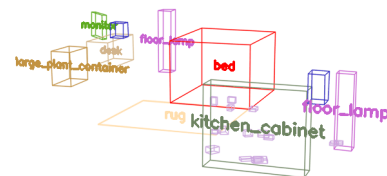


Figure 11. Our generated layouts exhibit strong adherence to the provided textual guidance.

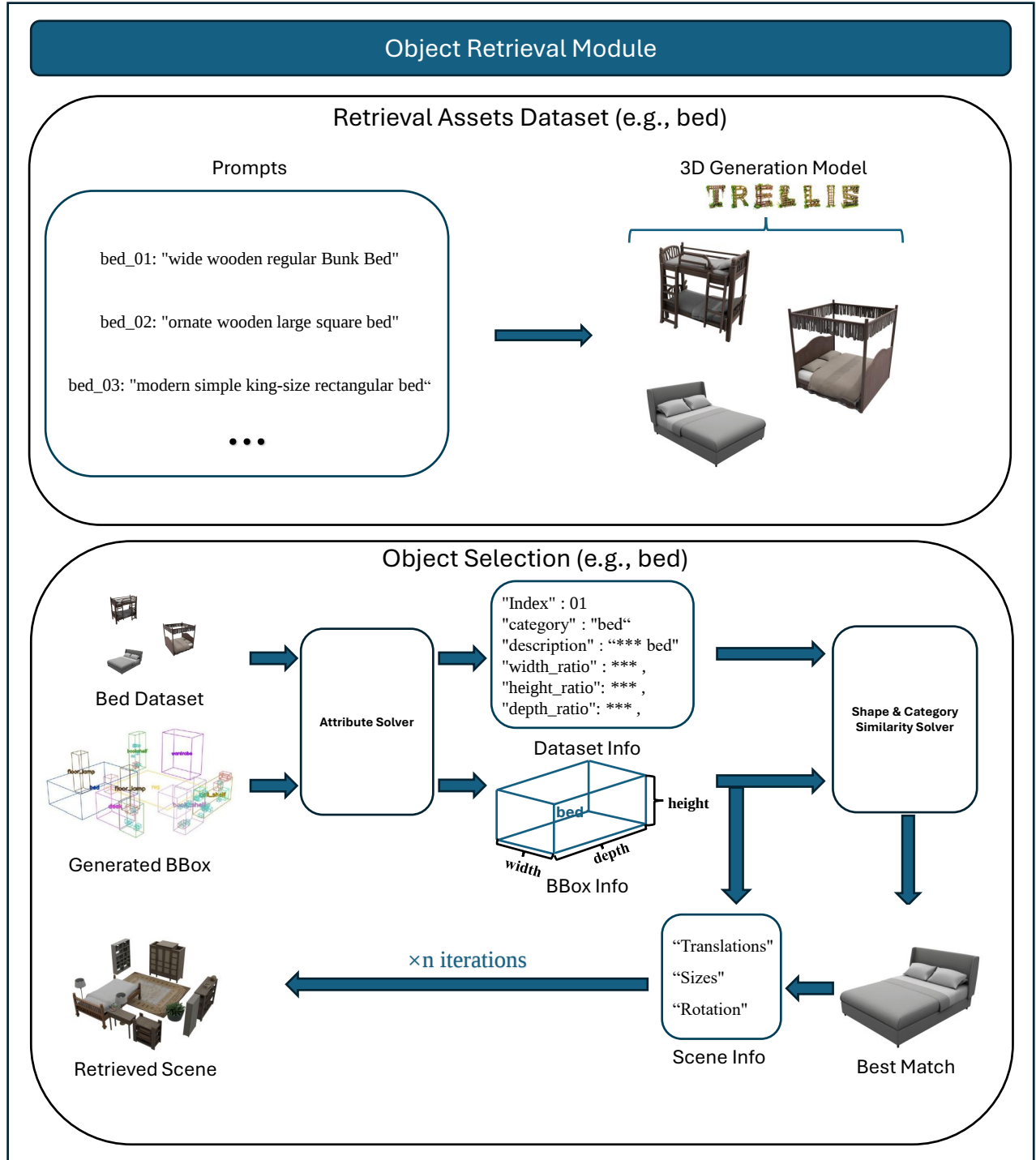


Figure 12. **Retrieval process of layout-to-scene.** This flowchart illustrates the process of retrieving and selecting 3D objects based on generated BBox information.

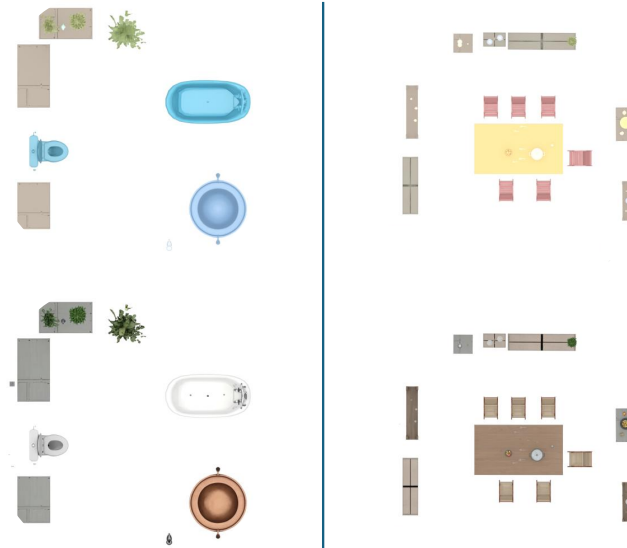


Figure 13. **Retrieval visualization of generated layouts.** The first row displays the retrieved 3D scenes' renderings with pure color schemes. The second row shows the retrieved 3D scenes' renderings with original textures applied.