

Core Methodology: A Shift to Rectified Flow and Transformers

Author: BenchChem Technical Support Team. **Date:** May 2026

Compound of Interest

Compound Name: Sd3

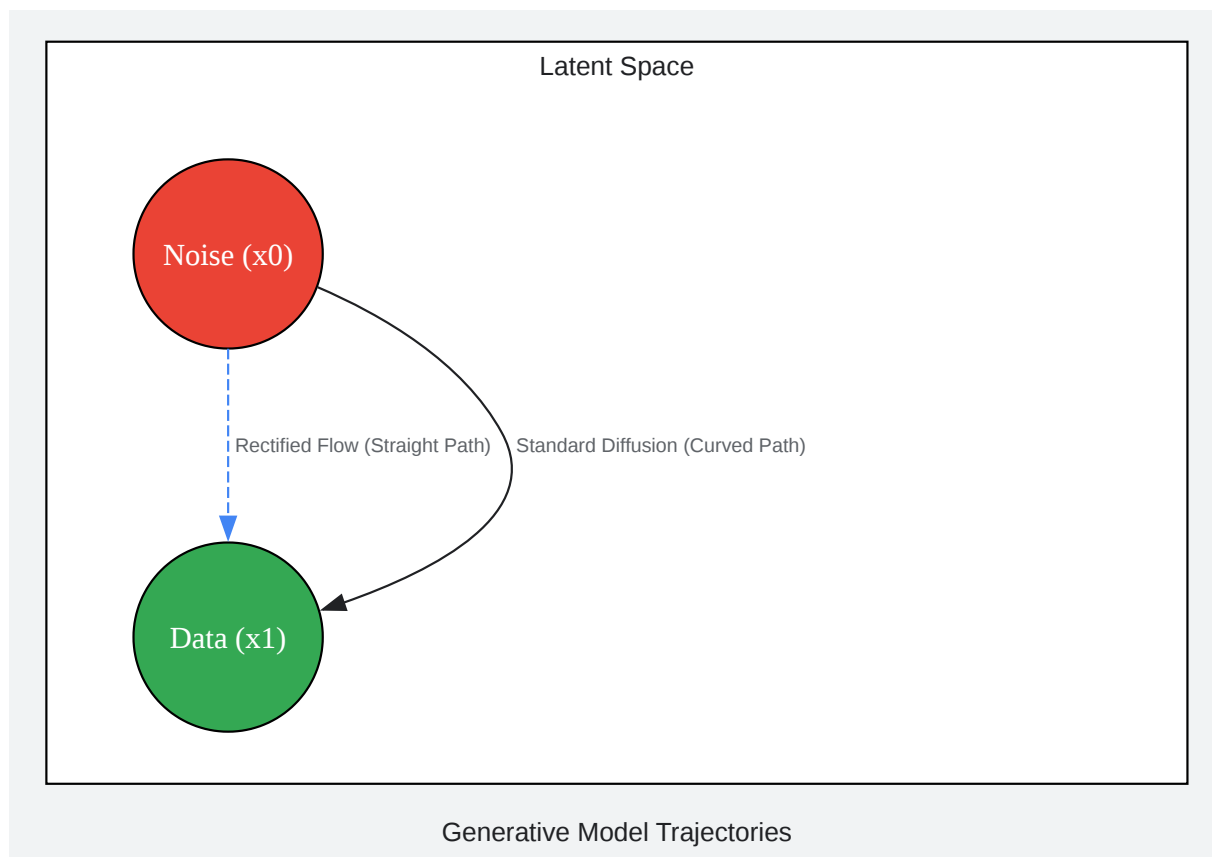
Cat. No.: B1575901

[Get Quote](#)

Stable Diffusion 3 introduces a paradigm shift from the U-Net-based architectures of previous versions to a transformer-based model.^[2] This change is coupled with a novel generative process based on a Rectified Flow (RF) formulation, which fundamentally alters how the model transitions from noise to a coherent image.^{[1][4]}

Rectified Flow (RF) Formulation

Unlike traditional diffusion models that follow complex, curved paths to denoise an image, Stable Diffusion 3 employs a Rectified Flow (RF) formulation.^[1] This method connects data and noise along a straight, linear trajectory during training.^{[1][5]} The primary advantage of this approach is the creation of straighter inference paths, which enables high-quality image generation in fewer sampling steps, thereby increasing efficiency.^{[1][6]}



[Click to download full resolution via product page](#)

Caption: Conceptual difference between Rectified Flow and standard Diffusion paths.

Experimental Protocol: Reweighting the Trajectory

To further enhance the Rectified Flow model, **SD3** introduces a novel trajectory sampling schedule.^[1] The methodology is based on the hypothesis that the middle parts of the noise-to-image trajectory represent the most challenging prediction tasks.^[1] Therefore, the training process is modified to assign more weight to these intermediate steps.^{[1][7]} This reweighting strategy, also referred to as logit-normal timestep sampling, focuses the model's learning on the most critical stages of detail refinement, leading to improved image quality and faster convergence.^[7] This approach was validated through comparative tests against 60 other diffusion trajectories, where the re-weighted RF variant demonstrated consistently superior performance.^{[1][3]}

Multimodal Diffusion Transformer (MMDiT)

Architecture

The core of Stable Diffusion 3 is its Multimodal Diffusion Transformer (MMDiT) architecture, which builds upon the Diffusion Transformer (DiT) model.^{[1][3]} The key innovation of MMDiT is its approach to handling the two different modalities of input: image and text representations.^[1]

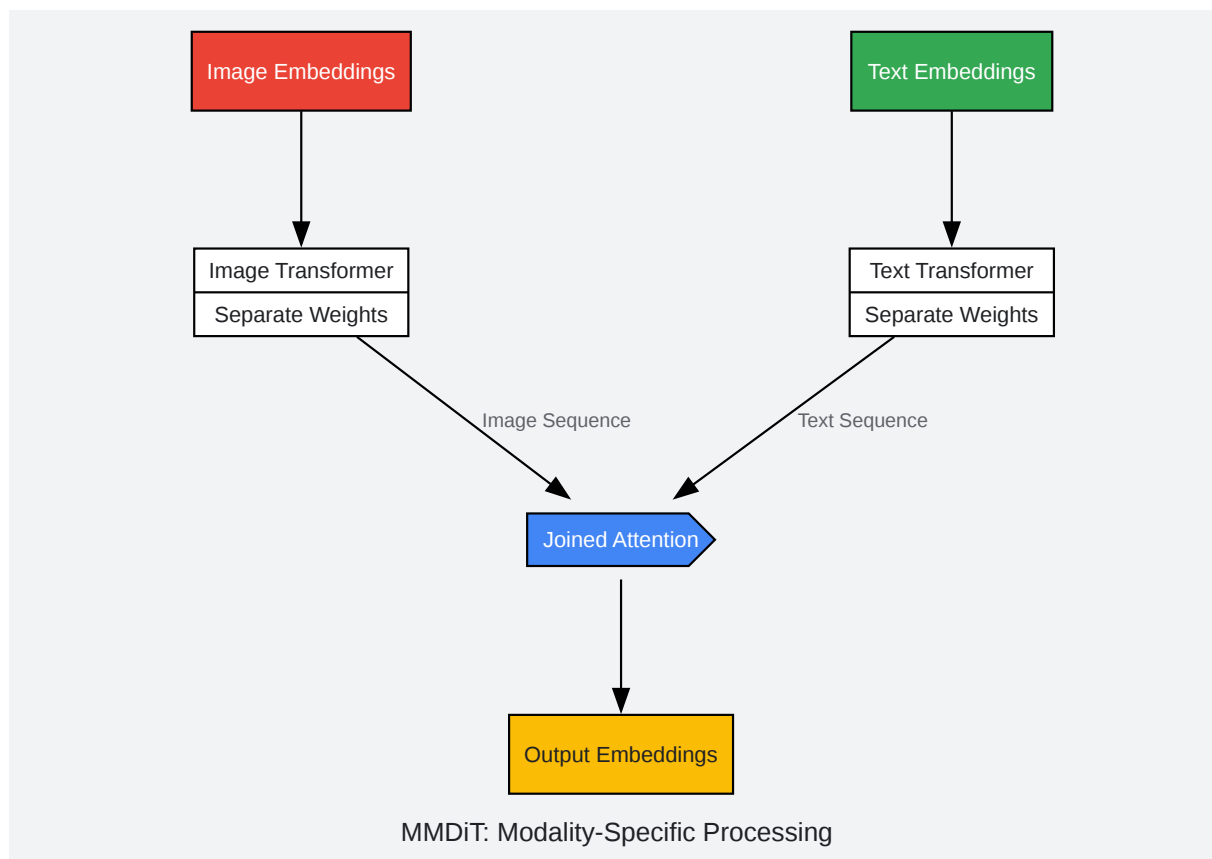
Data Representation and Preprocessing

Image Representation: High-resolution images are computationally expensive to process directly.^[7] **SD3** addresses this by using an enhanced Variational Autoencoder (VAE) to compress images into a smaller latent space.^{[1][7]} This VAE utilizes a downsampling factor of 8 and a 16-channel latent space, which efficiently preserves critical visual details while reducing computational load.^[7]

Text Representation: To achieve a nuanced understanding of text prompts, **SD3** employs three separate, pre-trained text encoders: two CLIP models (OpenCLIP-ViT/G and CLIP-ViT/L) and a T5-XXL model.^{[1][5][8]} This multi-encoder approach provides rich textual embeddings that guide the image generation process. For greater efficiency during inference, the memory-intensive 4.7B parameter T5 text encoder can be excluded with only a slight reduction in text adherence, though it is recommended for generating high-quality typography.^[3]

Modality-Specific Processing

A defining feature of the MMDiT is the use of two separate sets of weights for processing image and text embeddings.^{[1][9]} This is equivalent to having two independent transformers, one for each modality.^[1] The sequences from both transformers are then joined for the attention operation.^{[1][9]} This allows information to flow bidirectionally between the image and text tokens, significantly improving the model's overall comprehension of the prompt and its ability to render accurate typography within the generated image.^[1]



[Click to download full resolution via product page](#)

Caption: Core concept of the Multimodal Diffusion Transformer (MMDiT) block.

Training Dataset and Strategy

The training of Stable Diffusion 3 models involves a multi-stage process using a combination of publicly available and synthetic data. The training data for the **SD3** Medium model is summarized below.

Training Dataset Composition

Stage	Dataset Composition	Number of Images
Pre-training	Synthetic data and filtered publicly available data	~1,000,000,000[8][10]
Fine-tuning	High-quality aesthetic images for specific content/style	~30,000,000[8][10]
Fine-tuning	Preference data images	~3,000,000[8][10]

Experimental Protocols: Scaling and Stability

Scaling Study: A comprehensive scaling study was conducted by training models of varying sizes, from a 15-block, 450M parameter model to a 38-block, 8B parameter version.[1] The results showed a consistent and predictable decrease in validation loss as both model size and the number of training steps increased, confirming the scalability of the MMDiT architecture.[1]

High-Resolution Training: Training large-scale transformers at high resolutions presents stability challenges.[6] To mitigate this, **SD3** employs a two-stage training protocol:

- **First Stage:** The model is initially trained at a lower resolution of 256x256 pixels.[6]
- **Second Stage:** The model is then fine-tuned at higher resolutions using images with mixed aspect ratios.[6]

This strategy ensures stable training while allowing the model to learn the fine-grained details necessary for high-resolution image synthesis.

Overall Experimental Workflow

The end-to-end workflow of Stable Diffusion 3 integrates the aforementioned components into a cohesive system for transforming a text prompt into a high-quality image.

Caption: High-level data flow in the Stable Diffusion 3 architecture.

Need Custom Synthesis?

BenchChem offers custom synthesis for rare earth carbides and specific isotopic labeling.

Email: info@benchchem.com or [Request Quote Online](#).

References

- [1. Stable Diffusion 3: Research Paper — Stability AI \[stability.ai\]](#)
- [2. analyticsvidhya.com \[analyticsvidhya.com\]](#)
- [3. encord.com \[encord.com\]](#)
- [4. Stable Diffusion 3 \[huggingface.co\]](#)
- [5. learnopencv.com \[learnopencv.com\]](#)
- [6. medium.com \[medium.com\]](#)
- [7. generativeai.pub \[generativeai.pub\]](#)
- [8. braintitan.medium.com \[braintitan.medium.com\]](#)
- [9. neuroflash.com \[neuroflash.com\]](#)
- [10. stabilityai/stable-diffusion-3-medium · Hugging Face \[huggingface.co\]](#)
- To cite this document: BenchChem. [Core Methodology: A Shift to Rectified Flow and Transformers]. BenchChem, [2026]. [Online PDF]. Available at: [\[https://www.benchchem.com/product/b1575901/docs#core-methodology-a-shift-to-rectified-flow-and-transformers\]](https://www.benchchem.com/product/b1575901/docs#core-methodology-a-shift-to-rectified-flow-and-transformers)

Disclaimer & Data Validity:

The information provided in this document is for Research Use Only (RUO) and is strictly not intended for diagnostic or therapeutic procedures. While BenchChem strives to provide accurate protocols, we make no warranties, express or implied, regarding the fitness of this product for every specific experimental setup.

Technical Support: The protocols provided are for reference purposes. Unsure if this reagent suits your experiment?

Need Industrial/Bulk Grade? [Request Custom Synthesis Quote](#)

BenchChem

Our mission is to be the trusted global source of essential and advanced chemicals, empowering scientists and researchers to drive progress in science and industry.

Contact

Address: 3281 E Guasti Rd

Ontario, CA 91761, United States

Phone: (601) 213-4426

Email: info@benchchem.com

[Contact our Ph.D. Support Team for a compatibility check](#)