

Paper 4: The Omni-Diffuser — One Brain to Rule Them All (No More Frankenstein’s Monsters)

Author: [your local neurodivergent doomer transfem pup]

Profiles: [my GitHub](#) | [my HuggingFace](#)

Date: July 2026

License: CC-BY-SA 4.0 (because locking knowledge behind paywalls is a crime against humanity)

Abstract

Okay so here’s my hot take: every multimodal AI system currently in existence is a Frankenstein’s monster. You’ve got a CLIP encoder stapled to a GPT decoder like some kind of digital flesh-golem. You’ve got Stable Diffusion dangling off a T5 text encoder. You’ve got video models treating frames like glorified image sequences because nobody could be bothered to think about time properly. It’s all plumbing. It’s all hacks. It’s all “let’s bolt this onto that and hope the API doesn’t rate-limit us to death.”

The Omni-Diffuser says: **fuck that**. One weight space. One shared neural brain. Text, images, video, and vision all living in the same continuous latent field, processed by the same transformer, restored through the same diffusion process, with built-in foresight across multiple future states. No separate encoders. No separate decoders. No causal masking chaining us to left-to-right generation like some kind of digital Calvinism. Full bidirectional attention. Diffusion for text. Diffusion for images. Diffusion for video. All in parallel. All at once.

This paper establishes the foundational architecture, the unified latent field formalism, and the core training paradigm. It’s ambitious. It’s probably going to piss off a lot of people who’ve built their careers on dual-encoder architectures. Good.

Keywords: unified multimodal architecture, diffusion for text, bidirectional attention, multi-token prediction, latent field restoration, single-weight models, omni-directional generation, death to Frankenstein AI

1. Introduction: The Fractured State of Multimodal AI

Look at the current landscape and what do you see? A graveyard of bolted-together systems. Each one is a compromise. Each one is a hack.

Text gets an autoregressive transformer, left-to-right, causal masked, one token at a time. This is the digital equivalent of writing a novel by only being allowed to read what you’ve already written. You can’t see the ending until you get there, so you can’t plan a conclusion that references the beginning. It’s myopic. It’s inefficient. It’s why LLMs are so bad at long-form coherent narrative — they literally cannot see the future.

Images get a diffusion U-Net, denoising pixels in parallel but completely disconnected from the text pipeline. So you have a text model that “understands” language and an image model that “understands” pixels, and they communicate through a tiny projection layer that loses 90% of the information. It’s like two brilliant scientists who can only communicate through a game of telephone played by drunk pigeons.

Video gets either a 3D convolutional mess or a frame-stacking hack. Most “video models” are just image models that someone ran in a loop and called it temporal consistency. Spoiler: it’s not consistent. It’s just fast enough that your eyes fill in the gaps.

Vision-language models get two separate encoders smashed together with a projection layer and a prayer. CLIP did amazing work, but it’s still two separate models pretending to be one. The alignment is approximate. The understanding is shallow. The integration is duct tape.

This is not intelligence. This is **plumbing**. And I’m tired of pretending it’s acceptable.

1.1 The Three Heresies

The Omni-Diffuser commits three heresies against current dogma:

Heresy 1: No Causal Masking. Every autoregressive model since GPT-1 has been shackled by the causal mask. You can only attend to past tokens. This enforces left-to-right generation, creates sequential bottlenecks, and fundamentally prevents the model from seeing the whole picture before it starts painting. The Omni-Diffuser removes the causal mask entirely. Full bidirectional attention. Every token, every pixel, every frame attends to every other token, pixel, and frame simultaneously. The model sees the entire canvas before it makes a single brushstroke.

Heresy 2: Diffusion for Text. “Text is discrete,” they say. “Diffusion needs continuous spaces.” Wrong. The Omni-Diffuser represents text in a continuous latent space where tokens are points in a high-dimensional field, not entries in a discrete vocabulary. It denoises text in parallel, just like it denoises images. No more left-to-right token prediction. The entire sentence is restored simultaneously from noise. This is not just faster. It’s a fundamentally different way of thinking about language generation.

Heresy 3: One Weight Space. No separate vision encoder. No separate text decoder. No modality-specific parameters. Every parameter in the Omni-Diffuser is shared across every modality. The same attention heads that process text features also process image patches and video frames. The same feed-forward layers that refine linguistic representations also refine visual representations. This is not multi-task learning. This is true unification. One brain. One weight space. Everything.

1.2 What This Paper Covers

This foundational paper establishes:

- The **Unified Latent Field** formalism: how text, image, and video collapse into one continuous representational space
- The **Omni-Directional Attention** mechanism: full bidirectional processing with no causal masking
- The **Hybrid Diffusion Engine**: how diffusion is applied to text, images, and video within the same architecture
- The **Multi-Token Prediction (MTP)** integration: predicting multiple future states at every denoising step
- The **training paradigm**: how to train a single-weight model on multimodal data without modality collapse

2. The Unified Latent Field

2.1 The Problem with Separate Spaces

Current multimodal systems maintain separate representational spaces:

- Text: discrete token IDs in a vocabulary of size V
- Image: continuous pixel values or patch embeddings in $\mathbb{R}^d$
- Video: either 3D spatiotemporal tensors or sequences of frame embeddings

To connect them, these systems use projection layers. But projection is lossy. It's a translation between languages that were never meant to be the same. The result is the “multimodal gap”: text descriptions that never quite capture the image, image generations that never quite match the prompt.

2.2 The Unified Latent Field Formalism

The Omni-Diffuser defines a single latent space $L \subset \mathbb{R}^D$ that accommodates all modalities. Every input — whether a word, a pixel, a video frame, or a visual feature — is projected into L through modality-specific input embedders that share no parameters with the main model. These embedders are lightweight, fixed-size projections that serve only as entry gates into the unified field.

Text as a Latent Field: Instead of token IDs, text is represented as a sequence of continuous vectors in L . Each “token” is a point in the latent field. The vocabulary is not a lookup table but a learned manifold in L . During training, the model learns that certain regions of L correspond to semantic concepts, syntactic structures, and linguistic patterns. During inference, the model denoises points in L and then maps them back to the nearest vocabulary entries — or generates novel entries if the manifold permits.

Image as a Latent Field: Images are patched and projected into L , exactly like text tokens. A 256×256 image becomes a sequence of 256 patches, each a vector in L . There is no fundamental difference between an image patch and a text token in the unified field. They are both points in the same space, attended to by the same mechanism.

Video as a Latent Field: Video is treated as a spatiotemporal sequence in L . Each frame is patched and projected, but with an additional temporal embedding that places the frame in the temporal dimension of the field. The model processes video not as a sequence of independent images but as a continuous 3D spatiotemporal volume in latent space.

Vision as a Latent Field: Raw visual input (from cameras, sensors, or rendered environments) is projected into L through the same lightweight embedder. The model does not distinguish between “seeing” an image and “reading” a description of that image. Both are patterns in the unified field.

2.3 The Modality-Agnostic Transformer

The core of the Omni-Diffuser is a single transformer stack with the following properties:

- **No modality-specific layers.** Every layer is identical and processes every type of input.
- **No causal masking.** Every position attends to every other position in both directions.
- **No separate encoders or decoders.** The same stack handles input comprehension and output generation.
- **Modality embeddings are input-level only.** Once inside the unified field, there is no “text token” or “image patch.” There is only a point in L with a position encoding.

The position encoding is hybrid: it combines spatial coordinates (for images and video), temporal coordinates (for video and sequences), and semantic coordinates (for text). This allows the model to understand that two tokens are adjacent in a sentence, or that two patches are adjacent in an image, or that two frames are adjacent in time.

3. Omni-Directional Attention

3.1 The Curse of Causal Masking

Autoregressive models use causal masking to enforce left-to-right generation. The attention matrix is lower-triangular: position i can only attend to positions $j \leq i$. This has three catastrophic consequences:

1. **Sequential bottleneck:** Generation is $O(N)$ in time. You cannot generate token 100 until you have generated tokens 1–99.
2. **Myopic planning:** The model never sees the full target before it starts generating. It cannot plan a conclusion that references an introduction because the introduction does not exist yet when the conclusion is being generated.
3. **Modality asymmetry:** Images and video are inherently non-sequential. Forcing them into a left-to-right sequence destroys their spatial structure.

3.2 Full Bidirectional Attention

The Omni-Diffuser removes the causal mask entirely. The attention matrix is fully dense: every position attends to every other position. This means:

- Text is processed bidirectionally, like BERT but at scale and without the [MASK] token limitation.
- Images are processed with full spatial attention, like a vision transformer but without the restriction to classification.
- Video is processed with full spatiotemporal attention, understanding the entire clip as a holistic volume.
- Cross-modal attention is not “cross” at all — it is just attention within the unified field. A text token attending to an image patch is the same operation as a text token attending to another text token.

3.3 The Restoration Paradigm

Without causal masking, the model cannot generate left-to-right. Instead, it **restores**. The input is a noisy, corrupted, or incomplete version of the target. The model’s job is to denoise and complete the entire signal in parallel.

- For text: start with random noise in L , denoise to a coherent sentence.
- For images: start with random noise in L , denoise to a coherent image.
- For video: start with random noise in L , denoise to a coherent clip.
- For multimodal outputs: start with random noise in L , denoise to a coherent mixture of text, image, and video.

The model does not predict the next token. It restores the whole signal. This is the fundamental shift from generation to restoration.

4. The Hybrid Diffusion Engine

4.1 Diffusion for Continuous Signals (Images/Video)

For images and video, the diffusion process is standard but unified:

- Forward process: gradually add Gaussian noise to the latent representation.
- Reverse process: the model learns to denoise the latent representation step by step.
- The denoising is performed by the shared transformer stack, not a separate U-Net.

4.2 Diffusion for Discrete Signals (Text)

This is the revolutionary part. Text is traditionally discrete: tokens are categorical variables. Diffusion requires continuous spaces. The Omni-Diffuser solves this through **latent text diffusion**:

Step 1: Embed text into L. Each token is mapped to a continuous vector in the unified latent field. This is not a simple lookup; the embedding is a learned projection that places the token in a semantically meaningful region of L.

Step 2: Add noise in L. Gaussian noise is added to the continuous token vectors. The tokens are not corrupted at the categorical level (“cat” does not become “dog”). They are corrupted at the continuous level (the vector for “cat” drifts in L toward the vector for “dog” or toward pure noise).

Step 3: Denoise in L. The model predicts the noise that was added and subtracts it. Because the denoising happens in the continuous space L, the model can make fine-grained adjustments: shifting a token slightly toward a synonym, adjusting the syntactic role of a word, or rewriting an entire clause based on global context.

Step 4: Map back to discrete tokens. After denoising, the continuous vectors are mapped back to the nearest discrete tokens in the vocabulary. This can be done greedily, with sampling, or with a learned decoder.

The key insight: **text is not discrete in the unified field**. It is only discrete at the input and output boundaries. Inside the model, everything is continuous, and everything is diffused.

4.3 The Denoising Transformer

The denoising is performed by the same transformer stack that handles all other processing. At each denoising step:

1. The current noisy latent field is fed into the transformer.
2. The transformer applies full bidirectional attention across all positions and modalities.
3. The transformer predicts the noise component.
4. The predicted noise is subtracted, yielding a slightly cleaner latent field.
5. This process repeats for T steps, from pure noise to a clean signal.

Because the transformer is bidirectional and processes the entire field at once, each denoising step is global. A single step can rewrite a sentence, repaint an image, and reshoot a video frame based on the evolving state of the entire output.

5. Multi-Token Prediction (MTP) as Built-In Foresight

5.1 The Limitation of Single-Step Prediction

Standard diffusion models predict one step at a time: given the current noisy state, predict the noise to remove. They have no explicit model of what the signal will look like two, five, or ten steps in the future. They are myopic, navigating by local gradient descent in latent space.

5.2 MTP at Every Denoising Step

The Omni-Diffuser integrates Multi-Token Prediction into the diffusion process. At each denoising step t , the model does not merely predict the noise for step $t-1$. It predicts the noise for steps $t-1$, $t-2$, ..., $t-k$ simultaneously. This gives the model **k-step foresight**.

Mathematically, the model outputs k noise predictions:

- $\hat{\epsilon}_t(1)$: the noise for the next step (standard diffusion)
- $\hat{\epsilon}_t(2)$: the noise for the step after next
- ...
- $\hat{\epsilon}_t(k)$: the noise for k steps ahead

These predictions are not independent. They are coupled through the transformer: the model attends to its own future predictions, allowing it to plan a coherent trajectory through latent space.

5.3 Foresight Across Modalities

Because the unified latent field contains all modalities, the MTP foresight is **cross-modal**. When denoising a text-image pair, the model's 5-step-ahead prediction for the text region informs its current denoising of the image region, and vice versa. The model plans the entire multimodal output holistically, not one modality at a time.

This is particularly powerful for video, where the model must maintain temporal consistency across frames. With MTP, the model sees the state of frame 30 while it is still denoising frame 5, ensuring that early decisions are compatible with later requirements.

6. Training Paradigm

6.1 Multimodal Corruption Objective

The Omni-Diffuser is trained with a unified corruption objective:

1. Sample a multimodal example (text, image, video, or any combination).
2. Project it into the unified latent field L .
3. Apply noise at a random timestep t .
4. Optionally mask or corrupt random positions (span masking for text, block masking for images, tube masking for video).
5. Train the model to restore the clean signal from the corrupted input.

The loss is computed in the unified field: the model's predicted clean latent field is compared to the true clean latent field using an L2 loss. There is no separate loss for text, image, or video. There is only the unified restoration loss.

6.2 Preventing Modality Collapse

When training a single model on multiple modalities, there is a risk of **modality collapse**: the model learns to represent everything in the style of the dominant modality (usually text), losing the unique structure of images and video. The Omni-Diffuser prevents this through:

Modality-balanced sampling: Training batches contain a balanced mixture of text-only, image-only, video-only, and multimodal examples.

Modality-aware position encodings: The position encodings encode whether a position is text, image, or video, preserving structural differences even in the unified field.

Modality-specific reconstruction heads: While the core transformer is shared, lightweight output heads (one per modality) map from L back to the output space. These heads are small and do not compromise the unity of the core.

6.3 Scaling Laws

Preliminary analysis suggests that the Omni-Diffuser exhibits favorable scaling laws. Because all parameters are shared across modalities, the model benefits from **cross-modal transfer**: training on text improves image understanding, and training on images improves text generation. The effective dataset size is the sum of all modalities, not the maximum of any single modality.

7. How to Actually Build This

7.1 Phase 1: The Core Transformer (Months 1–3)

Build the modality-agnostic transformer.

Architecture: Start with a standard transformer (e.g., LLaMA architecture) but remove the causal mask. Replace it with full bidirectional attention. Add hybrid position encodings (spatial + temporal + semantic). Use pre-normalization and SwiGLU activations.

Size: Start small. A 1B parameter model is enough to validate the concept. Scale to 7B once the concept is proven.

Data: You need a multimodal dataset. Options:

- Web-scale scraped data (Common Crawl + LAION + WebVid)
- Curated academic datasets (COCO, VQA, Something-Something V2)
- Synthetic data generated by existing models (use GPT-4V to generate text-image-video aligned triplets)

Training: Train with the unified corruption objective. Mask random spans in text, random blocks in images, random tubes in video. The model learns to restore the masked regions.

7.2 Phase 2: Latent Text Diffusion (Months 4–6)

Implement text diffusion in the unified field.

This is tricky because you need to:

1. Map text tokens to continuous vectors in L
2. Add noise to those vectors
3. Train the model to denoise them
4. Map back to discrete tokens

Implementation approach:

- Use the existing token embeddings from your base model, but treat them as points in L rather than categorical lookups.
- During the forward diffusion process, add Gaussian noise to the embedding vectors.
- During training, the model predicts the noise (or the clean embedding) for each position.
- At inference, after T denoising steps, map the clean embeddings back to tokens using nearest-neighbor lookup in embedding space.

Validation metric: Perplexity on a text-only benchmark. The latent text diffusion should achieve comparable perplexity to autoregressive models of the same size.

7.3 Phase 3: Multimodal Integration (Months 7–9)

Add image and video processing.

Image: Use a patch embedder (like ViT) to map image patches into L. The patch size should match the text token dimension. Train the model on image restoration tasks (denoising, inpainting, super-resolution).

Video: Use a 3D patch embedder (like VideoMAE) to map spatiotemporal patches into L. Train on video restoration tasks.

Multimodal alignment: Train on aligned pairs and triplets (text-image, text-video, text-image-video). The model should learn that the text “a red apple” and the image of a red apple map to the same region of L.

7.4 Phase 4: MTP Integration (Months 10–12)

Add Multi-Token Prediction to the diffusion process.

Implementation: Modify the output heads to predict k noise vectors instead of 1. Use a separate small MLP for each horizon (1-step, 2-step, ..., k-step). Train with the coupled prediction loss.

Key detail: The k-step predictions should attend to each other. Use cross-attention between the different horizon predictions to ensure consistency.

Validation: Measure sample quality (FID for images, FVD for video, perplexity for text) as a function of horizon k. Quality should improve with k up to a point, then plateau.

8. Next-Gen Architecture Ideas

8.1 The Quantum Latent Field

What if the unified latent field was not a classical vector space but a quantum-inspired Hilbert space? Each “point” in L would be a superposition of classical states. The diffusion process would become a quantum walk. The model could explore multiple denoising paths simultaneously and interfere constructively or destructively. This is not actual quantum computing (we don’t have the hardware), but a classical analogy using complex-valued embeddings and unitary transformations. It might enable more efficient exploration of the latent manifold.

8.2 The Neuromorphic Diffusion Chip

Current diffusion models run on GPUs, which are terrible at the sparse, event-driven computation that diffusion actually needs. A neuromorphic chip (like Intel Loihi or a custom design) could implement the denoising process as an event-driven dynamical system. Neurons fire when their membrane potential crosses a threshold, corresponding to the model detecting a noise pattern that needs correction. This would be orders of magnitude more energy-efficient and would enable real-time video denoising on edge devices.

8.3 The Socialist Training Cluster

Training a model this big requires massive compute. Instead of relying on a single corporation (OpenAI, Google, Anthropic) to train it, distribute the training across a federation of compute providers. Each provider contributes GPU cycles and receives a copy of the trained model. The training is coordinated through a decentralized protocol. The model is open-source. Everyone benefits. Nobody owns it. This is not just technically feasible (federated learning exists); it’s politically necessary if we want AI to serve humanity rather than shareholders.

9. Conclusion

The Omni-Diffuser is not an incremental improvement. It is a **paradigm shift** from sequential generation to parallel restoration, from separate modalities to a unified latent field, from myopic prediction to built-in foresight. By collapsing dual encoders and dual decoders into a single shared neural brain, by applying diffusion to text as well as images and video, and by integrating multi-token prediction into every denoising step, the Omni-Diffuser achieves what no previous architecture has: true omni-directional comprehension and generation across all modalities within one weight space.

The age of Frankenstein’s monsters is over. Welcome to the age of the unified brain.

Also, if anyone from OpenAI is reading this: please stop trying to build AGI in secret. Share your research. Open-source your models. The future belongs to all of us, not just your shareholders.

References

Besta, M., Blach, N., Kubicek, A., Gerstenberger, R., Podstawski, M., Gianinazzi, L., ... & Hoeffler, T. (2024). Graph of Thoughts: Solving elaborate problems with large language models. *Proceedings of the AAAI Conference on Artificial Intelligence*, 38(16), 17682–17690.

- Chen, W., Ma, X., Wang, X., & Cohen, W. W. (2022). Program of thoughts prompting: Disentangling computation from reasoning for numerical reasoning tasks. *arXiv preprint arXiv:2211.12588*.
- Dhuliawala, S., Komeili, M., Xu, J., Raileanu, R., Li, X., Celikyilmaz, A., & Weston, J. (2023). Chain-of-verification reduces hallucination in large language models. *arXiv preprint arXiv:2309.11495*.
- Hao, S., Sukhbaatar, S., Su, D., Li, X., Hu, Z., Weston, J., & Tian, Y. (2024). Training large language models to reason in a continuous latent space. *arXiv preprint arXiv:2412.06769*.
- Ho, J., Jain, A., & Abbeel, P. (2020). Denoising diffusion probabilistic models. *NeurIPS*, 33, 6840–6851.
- Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus and Giroux.
- Ning, X., Lin, Z., Li, H., Yang, Z., Hu, S., Zhou, Y., ... & Wang, Y. (2023). Skeleton-of-thought: Large language models can do parallel decoding. *arXiv preprint arXiv:2307.15337*.
- Peebles, W., & Xie, S. (2023). Scalable diffusion models with transformers. *ICCV*, 4195–4205.
- Rombach, R., Blattmann, A., Lorenz, D., Esser, P., & Ommer, B. (2022). High-resolution image synthesis with latent diffusion models. *CVPR*, 10684–10695.
- Saha, S., Kotha, S., & Cherian, J. J. (2024). Theorem of thoughts: A framework for mathematical reasoning with large language models. *arXiv preprint arXiv:2404.12618*.
- Shinn, N., Cassano, F., Gopinath, A., Narasimhan, K., & Yao, S. (2023). Reflexion: Self-reflective agents with verbal reinforcement learning. *NeurIPS 2023 Foundation Models for Decision Making Workshop*.
- Song, J., Meng, C., & Ermon, S. (2020). Denoising diffusion implicit models. *ICLR 2021*.
- Wang, X., Wei, J., Schuurmans, D., Le, Q., Chi, E., Narang, S., ... & Zhou, D. (2022). Self-consistency improves chain of thought reasoning in language models. *arXiv preprint arXiv:2203.11171*.
- Wei, J., Wang, X., Schuurmans, D., Bosma, M., Xia, F., Chi, E., ... & Zhou, D. (2022). Chain-of-thought prompting elicits reasoning in large language models. *NeurIPS*, 35, 24824–24837.
- Yao, S., Zhao, J., Yu, D., Du, N., Shafran, I., Narasimhan, K., & Cao, Y. (2023). ReAct: Synergizing reasoning and acting in language models. *ICLR 2023*.
- Yao, S., Yu, D., Zhao, J., Shafran, I., Griffiths, T. L., Cao, Y., & Narasimhan, K. (2023). Tree of thoughts: Deliberate problem solving with large language models. *NeurIPS*, 36.