

Rendering the Future

How Krea AI Accelerates Foundation Model Training with WEKA

With WEKA® NeuralMesh™ as the data foundation for its multi-petabyte training cluster, Krea has achieved close to 100% GPU utilization on individual training workloads and approximately 40% Model FLOPs Utilization (MFU) across large distributed training runs.

Table of Contents

- 00.** Introduction
- 01.** The Challenge: A Storage System That Kept Breaking
- 02.** The Solution: WEKA NeuralMesh, a Foundation Worth Building On
- 03.** The Results: Creating Without Limits
- 04.** Looking Forward: The Picture Gets Bigger

Krea AI is an AI creative platform for visual professionals, aggregating 64+ image, video, and 3D AI models into one interface. Its proprietary foundation models, including Krea 2 for images and Realtime 14B for video, are built entirely in-house. Founded in 2022 in San Francisco, the company has grown to more than 30 million users and counts Shopify, Amazon Studios and Superside among its enterprise customers. It's privately held and backed by Andreessen Horowitz, Bain Capital Ventures and others.

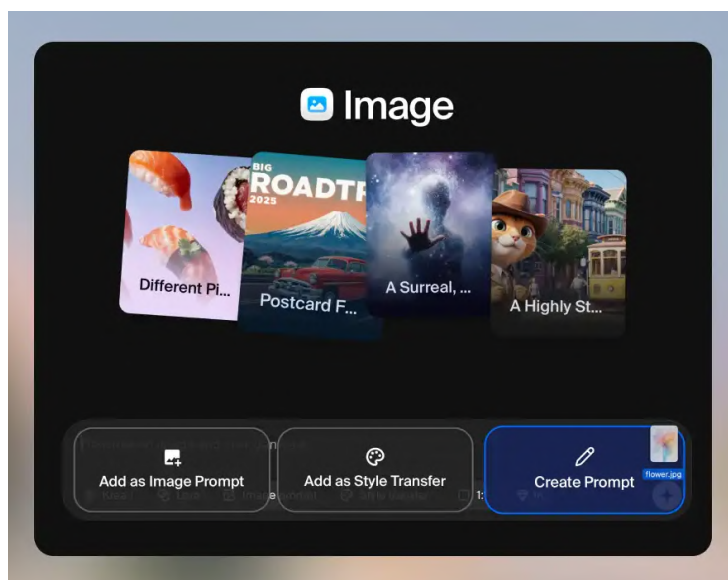
Driving that growth is a lean infrastructure team running hundreds of GPUs across a hybrid environment of bare-metal machines and rented GPU clusters on Kubernetes. Every model Krea ships depends on that infrastructure performing at scale. In a market where foundation model releases are measured in months, iteration speed is the competitive edge. Each hour a training run stalls is an hour a researcher can't test a hypothesis. Storage that couldn't keep pace was a direct constraint on what Krea could build.

The Challenge

A Storage System That Kept Breaking

As Krea scaled from model aggregation into frontier AI research, its storage requirements outpaced what open-source solutions could deliver. The team's workloads share a set of characteristics that stress-test distributed storage: large numbers of small files in training datasets, hundreds of concurrent clients reading and writing simultaneously, and checkpoint operations every 30 minutes during long pre-training runs. Ceph, their primary solution, couldn't handle the combination.

The failure modes were specific. Ceph struggled with directory-heavy workloads with many small files, causing severe metadata performance degradation at scale. It also lacked the fault tolerance the team needed: a node failure could bring the entire cluster down. When it wasn't failing outright, performance was highly inconsistent, forcing the team to constantly hand-tune folder structures and shard directories to keep it marginally functional.



Sangwu Lee, Krea's head of AI, admits, "Even how bad it was kept changing."

The operational cost was measurable. When Ceph failed, GPUs sat idle while the infrastructure team worked to restore service. Researchers couldn't distinguish storage failures from code failures, adding debugging overhead to every stalled run.

Will Beddow, head of supercomputing and AI infrastructure at Krea, put it plainly: "When your storage system becomes something you need to think about constantly, it becomes a very large, very real pain point."

Krea infrastructure engineer Gabriel Menezes felt that acutely as the one responsible for keeping researchers productive. "Researchers can't work for hours," he says. "Sometimes they don't know if it's their code or the storage."

A range of open-source alternatives, including other distributed file systems and object-backed solutions, fared no better. After pushing through the training run for Krea 1, the team had a clear read on the situation: their storage solution was the ceiling on what they could build. Krea 2 would require a different foundation entirely.

The Solution

WEKA NeuralMesh, a Foundation Worth Building On

The team ran a bake-off of every distributed storage solution they could find. Their requirements were non-negotiable: POSIX compliance, so the solution behaved like a standard file system for researchers; reliable performance under high concurrency with hundreds of simultaneous clients; consistent throughput on small-file workloads; and fault tolerance that kept the cluster running through node failures. NeuralMesh was the only solution that met all of them. “We tried a bunch of options,” Lee says. “WEKA was the only one that satisfied all our requirements.”

The team ran a 30-day proof of concept (POC). Researchers adopted it so quickly that the POC cluster became a core piece of production infrastructure. Menezes had told his team not to rely on it during the evaluation. “They were like: WEKA is so good, I’m going to put all my data there,” he says. Beddow: “The proof of concept stuck around for almost a year.”

Crafting the Setup

Krea runs NeuralMesh over InfiniBand on a custom Kubernetes topology. InfiniBand configuration varies significantly across cluster providers, and getting it right required close collaboration with WEKA’s engineering team. Beddow and Menezes both describe working through exotic configuration issues at any hour, with WEKA engineers engaging at a technical depth that matched Krea’s own infrastructure team. “I felt very confident,” Menezes says, “because I did not feel left alone.”

WEKA’s support model also differed from every other vendor Krea works with. Rather than waiting for Krea to report a problem, WEKA’s team proactively flagged anomalies in the cluster. “There were occasions where WEKA came to us and said: we are seeing something weird in your cluster, are you sure it’s OK,” Lee says. “That is quite rare when working with a software company.”

The Results

Creating Without Limits

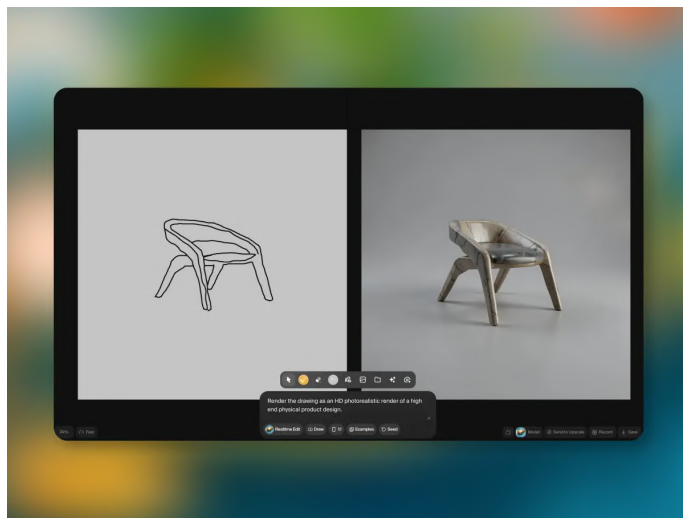
Krea's NeuralMesh cluster now spans several petabytes and serves hundreds of concurrent clients, handling pre-training runs, checkpointing, image processing and day-to-day research from a single storage layer. The performance numbers are a direct consequence of storage that no longer fails.

- **~100% GPU utilization on individual training workloads**, up from extended periods of idle time during Ceph failures when the team had to halt runs and restore the cluster.
- **Approximately 40% MFU across large distributed training runs**. By Lee's assessment, this is a number even well-resourced engineering teams rarely sustain. A good MFU for training large AI models typically ranges from 35% to 45%.
- **1 TB reads and 1 TB writes simultaneously**, sustained without performance degradation across training, checkpointing, and image processing workloads running in parallel.
- **Zero training failures attributable to storage since deployment**. "We have not had a single training fail due to the storage solution since we implemented WEKA," Beddow says.
- **Node-level fault tolerance**. If a node goes down, the cluster keeps running. Ceph didn't provide that.
- **Thousands of engineering hours saved** by not building a custom object-storage fallback. Model release timelines have accelerated by **weeks to months**.
- **Improved latencies and throughput on development machines**, giving researchers a storage layer that behaves like local SSD regardless of workload size.

Krea 1 and Krea 2, the company's first foundation image models, were both trained on NeuralMesh.

"Within just half a year, we trained a foundation model from scratch that is competitive by global standards," said Lee. Those models now serve millions of creatives across architecture, film, VFX and design.

Menezes describes watching Krea 2 train as the moment the difference became concrete. "We would checkpoint every 30 minutes and see the write throughputs going through the roof on the dashboard," he says. "WEKA would not break a sweat." Beddow on what NeuralMesh ultimately unlocked: "NeuralMesh made it possible to train models at the scale we needed."



Lee captures the operational shift in a single line: “Our storage system went from complete chaos to order.” The more personal version: “You take files and I/O for granted, but when you’ve had it taken away for a long time, and it’s given back, it’s like going camping and finally taking a shower. I will not take this for granted ever again.”

Looking Ahead

The Picture Gets Bigger

Krea expects its compute and storage environment to grow by 5 to 10 times in the next year. That trajectory means the team will need to manage significantly more GPUs, larger training clusters, and more concurrent clients and datasets than what it runs today. Every constraint that held them back with Ceph, including small-file performance, concurrency limits, and inconsistent throughput, would have compounded at that scale. NeuralMesh was built for it.

Beddow describes NeuralMesh as one of the few distributed storage solutions he’s used that works as advertised out of the box. “The difference between NeuralMesh and our earlier solutions is night and day.”

For a team with this much to build and this little room for distraction, storage that scales without demanding attention is the foundation everything else depends on.

