

Building a More Reliable Foundation for Research

How Oxford's Advanced Research Computing Service Modernized Its Storage with WEKA® NeuralMesh™

The University of Oxford's central HPC and HTC service replaces high-performance scratch storage to gain serviceability, fault isolation across shared infrastructure, and a single platform spanning InfiniBand and Ethernet.

Table of Contents

00. Introduction

01. The Challenge: Serviceability, Blast Radius, and Two Fabrics, One Namespace

02. The Solution: One Distributed Namespace, Native on Both Fabrics

03. The Impact: Less Firefighting, More Headroom

04. Looking Forward

The University of Oxford's Advanced Research Computing

(ARC) service runs the university's central high-performance and high-throughput computing infrastructure for non-sensitive research data, supporting researchers across every division, from genomics and protein structure prediction to computational chemistry, archaeology, and a growing volume of AI and large language model work.

ARC's current footprint includes between 15,000 and 17,000 compute cores, a couple hundred GPUs, and several petabytes of storage, all running within a 600-kilowatt power envelope. Two compute clusters sit behind that footprint, one on InfiniBand for tightly coupled, multi-node MPI and OpenMP workloads, and one on Ethernet, including newer 100-gigabit nodes, for high-throughput and GPU work. Around 1,500 researchers are drawing on these resources at any given time, with several thousand more cycling through over the course of a doctoral or postdoctoral project.

Dr Steven Young, Head of Technical Services in Oxford ARC, summed up what draws him to the work: “We get to play with the new toys, but what really matters is that we provide the expertise that enables researchers to access hardware they wouldn’t be able to procure within their own research group. We provide the scale for people to do larger, more significant work on this central resource. What gets me excited is that we’re helping researchers across the university achieve their research goals.”

The Challenge: Serviceability, Blast Radius, and Two Fabrics, One Namespace

As a high-performance computing service, ARC's central job has always been to deliver data to compute and GPU resources fast enough to keep them busy, what Dr. Young calls "feeding the beast." The scratch storage ARC ran before WEKA® NeuralMesh™ made the job harder than it needed to be.

"Whenever we needed to fix things or make things work, or things broke, there were headaches. Moving to upgrade or fixing things by installing new firmware involved varying amounts of pain," Dr. Young said. "We were looking for a solution that didn't introduce more pain points and didn't break things in different ways. **Serviceability** is maybe the best word for it. We wanted to get away from the pain of not being able to keep current with the manufacturer's recommendations."

Reliability carried a related but more significant concern. Yassamine Mather, Scientific Software Developer in Oxford ARC, recalled, "Before WEKA, our main concern was that a single component failure could propagate beyond the affected subsystem and compromise adjacent parts of the environment."

Dr. Young put it more directly: "We were worried about **a rogue job causing a problem to the storage**, which caused more than just a problem for that job, but a problem for everyone across the whole of our estate."

Capacity was a real but secondary factor, rather than a wall ARC was hitting. Asked directly whether the previous solution was limiting what researchers could do, Dr. Young was candid: "We knew we wanted to get faster storage, able to cope with more compute resources calling on it. That probably was a motivator, but not necessarily a wall."

The fabric question was the hard technical constraint with no easy workaround. ARC's two compute clusters run on different fabrics: InfiniBand for low-latency message passing and large-memory jobs spanning multiple nodes, and Ethernet for higher-throughput and GPU work. Any storage system ARC adopted had to **present a single namespace across both with full performance on each.**

The Solution: One Distributed Namespace, Native on Both Fabrics

ARC ran a competitive evaluation to replace its high-performance scratch storage, which WEKA won. As Dr. Young described it, "WEKA won the competition we ran to find out what high-performance storage we should go for. Part of the criteria was, **can it feed the data to our compute and GPU resources appropriately** running on the networking we have now and the networking we might have in the next two to five years."

WEKA NeuralMesh is a storage and memory software platform. It pools NVMe across the cluster into a single, high-performance namespace, with data and metadata distributed rather than concentrated on a small number of controllers. That architecture is the direct technical answer to the blast-radius problem Dr. Young described. Because no single node or client holds an outsized share of the system's state, a misbehaving job can no longer take the rest of the storage down with it. Dr. Young confirmed the outcome directly: "With NeuralMesh, we know that one client can't kill the whole storage system."

NeuralMesh also resolved ARC's fabric requirement. "Our storage needs to be able to work across both InfiniBand and Ethernet fabrics, which WEKA does easily for us," Dr. Young said. "It's able to present the system we have on NeuralMesh to both fabrics without us having to worry about it."

The deployment wasn't without complications. At installation, ARC's server hardware included two generations of network cards from the same manufacturer, and a software incompatibility between them, tied to the underlying DPDK networking libraries NeuralMesh's data path relies on, meant the two card types couldn't be used together as expected.

"WEKA came to the table and helped solve the problem for us," Dr. Young said. "We were very glad that the presales team stepped up and made our hardware work for us." That support has continued past deployment. "We have access to support through a Slack channel, which has been great," Dr. Young said.

Beyond replacing scratch storage outright, NeuralMesh allowed ARC to carve out persistent storage areas within the same namespace for datasets researchers use repeatedly. "We now provide persistent areas for some user datasets, like the AlphaFold databases, so people can use that rather than the scratch storage where they have to copy data up," Dr. Young said. "Otherwise they waste ten minutes of a day-long computational run just copying data up and pulling results back down."

The Impact: Less Firefighting, More Headroom

The result was a storage layer ARC no longer has to manage around.

- Storage capacity for high-performance work grew well beyond the previous setup. "The previous solution was at 50-100 TB scale, whereas now we're at much more."
- The blast-radius risk the team lived with before is gone. "With NeuralMesh, we know that one client can't kill the whole storage system," said Dr. Young.
- The team no longer has to think about storage day to day. "With WEKA, I know that the data I have is accessible where it needs to be. We don't have to worry about storage," said Dr. Young. "Mather agreed: 'It's a reliable resource.'"
- Support has become a resource the team reaches for directly rather than escalates through.

Looking Forward

Beyond the infrastructure itself, what motivates the team is watching it translate into results researchers can point to. Mather put it simply: “One of the good moments is when a research group says they’ve produced a paper in a very prestigious journal and named ARC as essential to their research. It’s a very good moment because we see that what we’ve done has made a difference to their publication.”

That's the foundation NeuralMesh supports: persistent storage and memory for high-demand datasets, a single distributed namespace that natively handles both of ARC's clusters, and a support relationship the team can lean on directly, giving ARC the headroom to keep helping researchers turn computational scale into results worth publishing.

See what NeuralMesh can do for your own research computing environment.

Talk to WEKA about running high-performance and AI research workloads on NeuralMesh.

Visit weka.io to get started.