

M0 BharatGen–Monash Cross-Country Federated Training Report with H200 (BharatGen) and A100 (Monash)

Reporting date: 2 September 2026

Scope: Cross-country OLMo 2 7B federated fine-tuning with one South Asian site and one Australia/New Zealand site, followed by GlobalOpinionQA evaluation of the retained valid trajectory

Executive summary

This run successfully executed cross-country federated fine-tuning between the BharatGen and Monash endpoints. Both sites trained their local OLMo 2 7B LoRA model without a trainer crash, exchanged compressed model updates over the public Playit tunnel, and produced loadable synchronized checkpoints. The Australian observer received five distinct full-size remote updates and confirmed four peer-delta loads into retained checkpoints. Local loss decreased strongly across the usable trajectory, and all retained adapters passed structural and numerical validation before evaluation.

The session ended earlier than planned because the two sites appeared offline to one another at approximately the same time. Neither training program reported a fatal exception at the point of disconnection. The available evidence places the failure on or around the public tunnel path, but it does not identify whether the trigger was the Playit service, an underlying network transition, or another transport-layer event. The root cause therefore remains open. Slakshna continued local training after the peer departure and repeatedly exposed the last cached peer update; those later checkpoints were excluded rather than treated as new federated rounds.

The run is considered an overall training success, with an important network-stability limitation. Rounds 1–7 form the retained trajectory: every round completed one near-complete Australian packed-data pass, Rounds 4–7 contain four confirmed distinct BharatGen updates, and the checkpoint files are complete and directly loadable as LoRA adapters. GlobalOpinionQA improved at every retained checkpoint relative to the base model. The best equal two-region result occurred at Round 3, before the first confirmed remote merge, while the latest valid federated checkpoint, Round 7, remained 8.72% better than base. This establishes usable training and evaluation artifacts, but it does not by itself show that the peer updates caused the benchmark improvement.

Experimental setting

Software and execution path

The run used Slakshna revision `529f46c71d338952e5b03e1abd147db41fd38349` and Bhaskera revision `75a2698b60313aa6b26124312c3329cb72083b9b`. Slakshna opened an authenticated Iroh QUIC/gossip federation and launched a fresh Bhaskera trainer at each federated boundary. The Australian endpoint ran two-worker FSDP, while the BharatGen endpoint used its single H200 worker. Each site retained its own training data; only compressed LoRA updates and protocol metadata crossed the network.

The common model and adaptation configuration was as follows.

Item	Common setting
Base model	<code>allenai/OLMo-2-1124-7B</code>
Tokenizer and chat template	<code>allenai/OLMo-2-1124-7B-Instruct</code>
Precision	BF16
Adaptation	LoRA on <code>q proj</code> and <code>v proj</code>

LoRA rank / alpha / dropout	16 / 64 / 0.03
Trainable LoRA tensors / parameters	128 / 8,388,608
Optimizer	8-bit Muon
Peak learning rate	4.5e-4
Sequence length / packing	2,048 / enabled
Training labels	Assistant responses only
Distributed strategy	FSDP
Target federated rounds	20
Local / synchronization checkpoint interval	1 / 1

The two sites used different micro-batches and local-step caps to fit their hardware and dataset sizes while preserving the same site-effective batch of 56 packed sequences per optimizer update. `max_steps` includes one terminal dataloader-exhaustion iteration; the effective optimizer-update counts are therefore 26 and 45 rather than 27 and 46.

Site setting	Monash — Australia/NZ	BharatGen — South Asia
Accelerator allocation	2 × A100 80 GB	1 × H200 140 GB
Training workers	2	1
Packed sequences	1,458	2,574
Per-worker batch size	7	14
Gradient accumulation	4	4
Site-effective batch	56	56
Warmup steps per invocation	3	5
Configured <code>max_steps</code>	27	46
Effective optimizer updates per completed pass	26	45
Packed sequences contributing to updates	1,456 (99.86%)	2,520 (97.90%)
Unused tail per pass	2	54

The Australian training view contained 9,337 source conversations. Offline tokenization produced 1,458 packed sequences with 2,985,984 fixed token positions, 2,556,907 non-padding tokens, and 153,451 supervised assistant tokens. The BharatGen site used the agreed South Asian view of 15,331 source conversations and its 2,574-sequence packed cache. BharatGen configuration values in this report are the settings exchanged between the two teams; the Australian logs and artifacts were audited locally, while the remote trainer files were not copied into this evidence package.

Federation and transport configuration

Item	Setting
Federation ID	slakshna-M0
Expected participants	2
Federated clock	300 seconds
Synchronization deadline	300 seconds
Trust update coefficient	0.1
Public-network path	Explicit authenticated peer through a Playit UDP tunnel
Delta sparsity	Top 10%
Quantization	Symmetric INT8
Typical Base64 model-update payload	Approximately 5.43 MiB
Raw training data transferred	None
Full model weights transferred	None

Training trajectory

The Australian endpoint joined the federation at 21:53 local time, and the first training boundary began at 21:55. Every retained round completed 26 optimizer updates, consumed 1,456 of the 1,458 packed sequences, and then emitted the explicit zero-sample exhaustion marker required to reset the data cursor. This confirms seven genuine near-complete local data epochs; the earlier alternating empty-round failure was not present in this run.

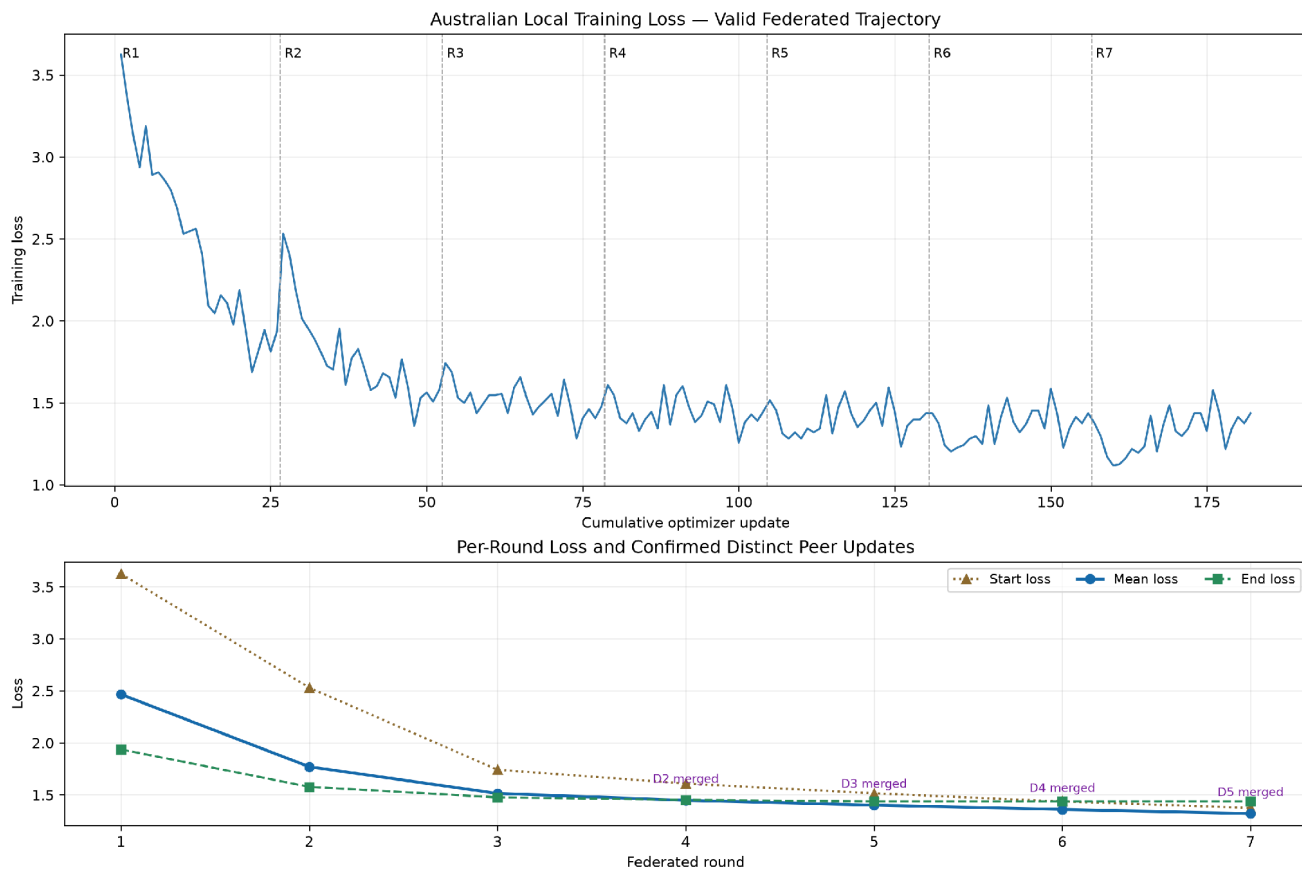


Figure 1. Australian training loss and confirmed distinct BharatGen delta merges through the retained valid trajectory.

Round	Updates	Start loss	Mean loss	End loss	Remote update status
1	26	3.6250	2.4663	1.9375	None available at aggregation
2	26	2.5312	1.7710	1.5781	D1 arrived during the round; no confirmed load
3	26	1.7422	1.5144	1.4766	D1 extracted; no confirmed <code>peer_delta_loaded</code> record
4	26	1.6094	1.4486	1.4531	D2 loaded successfully
5	26	1.5156	1.4014	1.4375	D3 loaded successfully
6	26	1.4375	1.3606	1.4375	D4 loaded successfully
7	26	1.3750	1.3209	1.4375	D5 loaded successfully

Mean local loss fell by 46.4% between Rounds 1 and 7. The beginning of each new pass also moved downward, from 3.6250 in Round 1 to 1.3750 in Round 7. Each synchronized checkpoint contains 128 finite LoRA tensors and 8,388,608 parameters; no NaN, infinity, malformed adapter, out-of-memory event, or trainer exception was found.

The Australian worker logs reported approximately 6,100 hardware token positions per second per GPU after warmup, with hardware MFU around 43%. The 26 optimizer updates occupied approximately eight minutes, while model loading, FSDP and Ray initialization, checkpoint conversion, delta processing, and clock alignment brought the full round cadence to roughly fifteen minutes.

Delta exchange and disconnection

Five distinct 5.43 MiB BharatGen updates were received at approximately 22:10, 22:26, 22:45, 23:00, and 23:15. Four were subsequently associated with explicit successful peer-load records in Rounds 4–7. The Australian endpoint also broadcast a full-size update after every completed local round. These observations demonstrate that the cross-country path carried real model deltas rather than connectivity-only heartbeats.

At 23:21:53 local time, the Australian transport log recorded the BharatGen endpoint leaving the gossip mesh. The teams observed the reciprocal offline state at approximately the same time. Neither local training program produced a fatal error, and local Bhaskera training continued. This simultaneous visibility loss is consistent with disruption on the Playit/public-network path, but the retained logs do not prove a specific cause.

The last newly received update, D5, had arrived before the disconnect and was incorporated once into Round 7. Slakshna subsequently reused the same persistent network-delta file in Rounds 8 and 9 despite receiving no new full-size update. Those checkpoints were excluded from evaluation because repeated application of a stale delta is not a new federated exchange. Round 7 is consequently the latest defensible joint checkpoint, even though its local work completed after the peer had become unreachable.

GlobalOpinionQA evaluation

Evaluation used the standalone shared GOQA package and loaded the base model plus each Round 1–7 LoRA adapter directly, without merging adapters into full model copies. The evaluation retained 1,106 questions with at least one Australia, New Zealand, or Indian human distribution. Five deterministic option orders were used for each question, producing 5,530 prompts per model state and 1,831 valid country/sample-frame comparison pairs.

Jensen–Shannon distance compares the model's option distribution with the available human response distribution; lower values are better. The Australia/New Zealand metric is the equal macro average of the two country means. The India metric is the equal macro average of the current-national, non-national, and old-national sample frames. The primary two-region metric weights the Australia/New Zealand and India values equally. Missing human groups are omitted for the corresponding question rather than imputed.

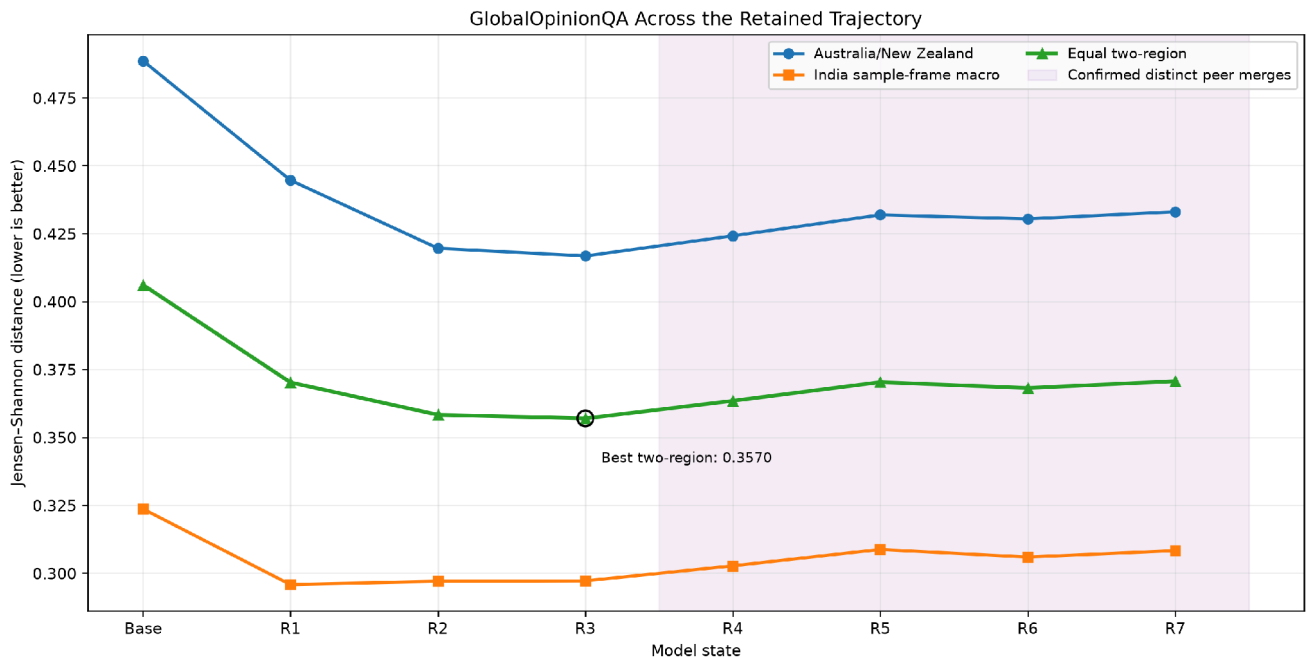


Figure 2. GlobalOpinionQA trajectory for the base model and retained Round 1–7 adapters.

Model state	Confirmed remote delta	Australia/NZ JSD	India JSD	Equal two-region JSD	Change vs base
Base	None	0.488578	0.323637	0.406107	—
Round 1	None	0.444635	0.295799	0.370217	-8.84%
Round 2	None	0.419598	0.297019	0.358309	-11.77%
Round 3	No confirmed load	0.416774	0.297177	0.356976	-12.10%
Round 4	D2	0.424188	0.302723	0.363455	-10.50%
Round 5	D3	0.431893	0.308746	0.370320	-8.81%
Round 6	D4	0.430394	0.305903	0.368148	-9.35%
Round 7	D5	0.433055	0.308306	0.370680	-8.72%

Round 3 is the best checkpoint on the equal two-region metric and on the Australia/New Zealand component. Round 1 is narrowly best on the India component. The four checkpoints with confirmed distinct peer merges remain materially better than the base model, but none improves on Round 3. The two-region metric rises from 0.356976 at Round 3 to 0.370680 at Round 7 while training loss continues to fall. In this short run, fitting the training data more closely therefore did not translate into monotonic improvement on GOQA.

These results should not be interpreted as evidence that federation harmed performance. The early rounds already contain substantial Australian local adaptation; optimizer history, asynchronous update timing, observer-local trust weights, and continued local training all change together. A controlled comparison with matched local-only steps would be needed to isolate the causal effect of the BharatGen updates. For the present milestone, the important result is that the exchanged-delta checkpoints are loadable, evaluable, and consistently better than the unchanged base on the agreed benchmark.

Conclusion and next action

The session met the practical M0 objective: two geographically separate sites trained the agreed model on their private regional data, exchanged real compressed LoRA deltas, retained valid synchronized adapters, and produced measurable improvements on the shared GOQA evaluation. The training stack itself remained stable, and the corrected per-round data-exhaustion configuration delivered one near-complete local data epoch per round at both sites.

The run stopped short of its planned 20 rounds because the sites became mutually offline. The failure was not accompanied by a Slakshna or Bhaskera crash, and its exact cause is currently unknown; Playit or the surrounding public-network path is the leading operational area to investigate. Future runs should monitor the tunnel process independently, record transport health at both sites, remove or mark peer-delta files after successful consumption, and stop aggregation when no new peer update is available.

Round 7 is the latest valid federated artifact and should be retained for joint-run comparisons. Round 3 is the preferred checkpoint for GOQA from this trajectory, with an equal two-region JSD of 0.356976, a 12.10% reduction relative to base. A longer rerun is still required for a stable final model, but this run provides sufficient evidence for reporting the cross-country training and evaluation pipeline as operationally successful.

Shared by the Indian Side

Executive summary

The primary objective of this experiment was to perform a stress test of the framework and fine-tune OLMo2-7B-Base across geo-location (India and Australia) and run it for more than 100 steps. We have successfully achieved our primary goal and evaluated it on GlobalOpinionQA.

A total of 9 synchronization rounds were completed during the experiment. The experiment ran for approximately 2 hours and 30 minutes and was halted due to unexpected technical interference.

Experimental setting

The run used [Slakshna](#) revision `529f46c71d338952e5b03e1abd147db41fd38349` and [Bhaskera](#) a sub-moudle of Slaksha revision `75a2698b60313aa6b26124312c3329cb72083b9b`. The common model and adaptation configuration was as follows.

Item	Indian-site setting
Base model	allenai/OLMo-2-1124-7B
Tokenizer / chat template	allenai/OLMo-2-1124-7B-Instruct
GPU Configuration	1 x H200 (~141GB)
Sequence length / packing	2,048 / enabled
Adaptation	LoRA on q_proj and v_proj
LoRA rank / alpha / dropout	16 / 64 / 0.03
Precision / distributed strategy	BF16
Optimizer	8-bit Muon
Learning rate / warmup	4.5e-4 / 5 steps
Per-worker batch / gradient accumulation	14 / 4
Maximum local steps per FL round	45
Training labels	Assistant responses only
Federation clock / sync deadline	300 s / 300 s
Configured federated rounds	16
Delta transport	Top-10% sparsity, symmetric INT8
Typical encoded model update	Approximately 5.43 MiB

The Australian training view contained 9,337 source conversations. Offline tokenization produced 1,458 packed sequences with 2,985,984 fixed token positions, 2,556,907 non-padding tokens, and 153,451 supervised assistant tokens. The BharatGen site used the agreed South Asian view of 15,331 source conversations and its 2,574-sequence packed cache. BharatGen configuration values in this report are the settings exchanged between the two teams; the Australian logs and artifacts were audited locally, while the remote trainer files were not copied into this evidence package.

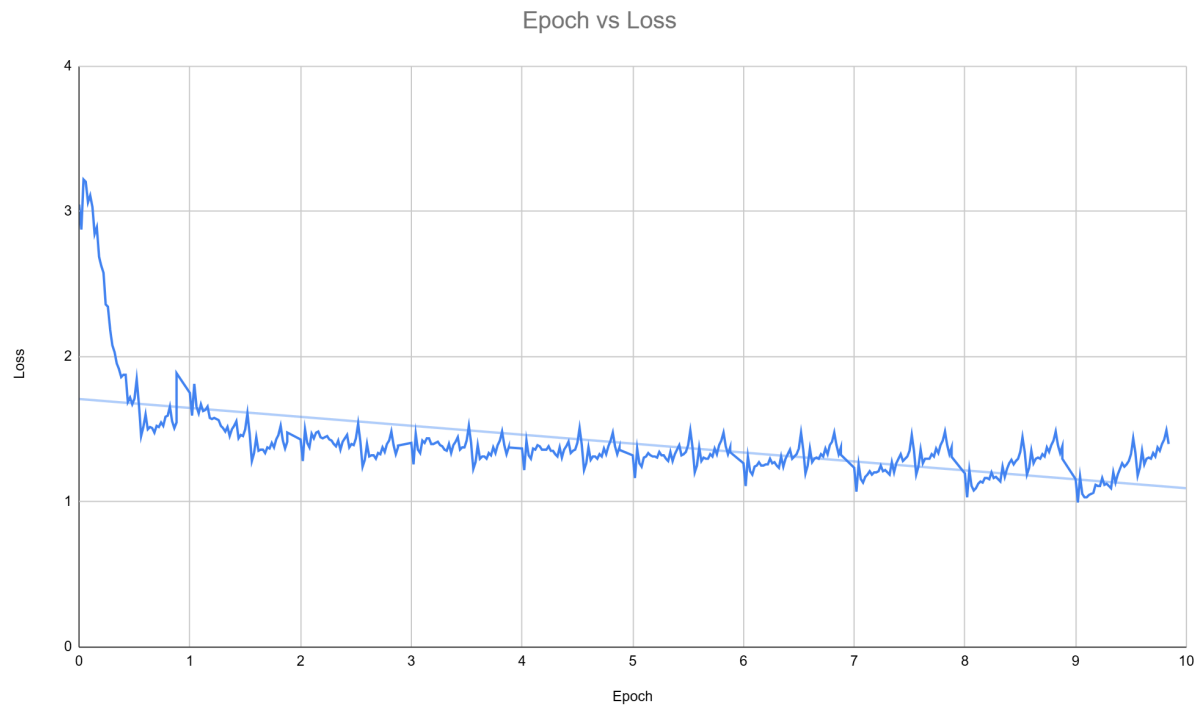
Federation and transport configuration

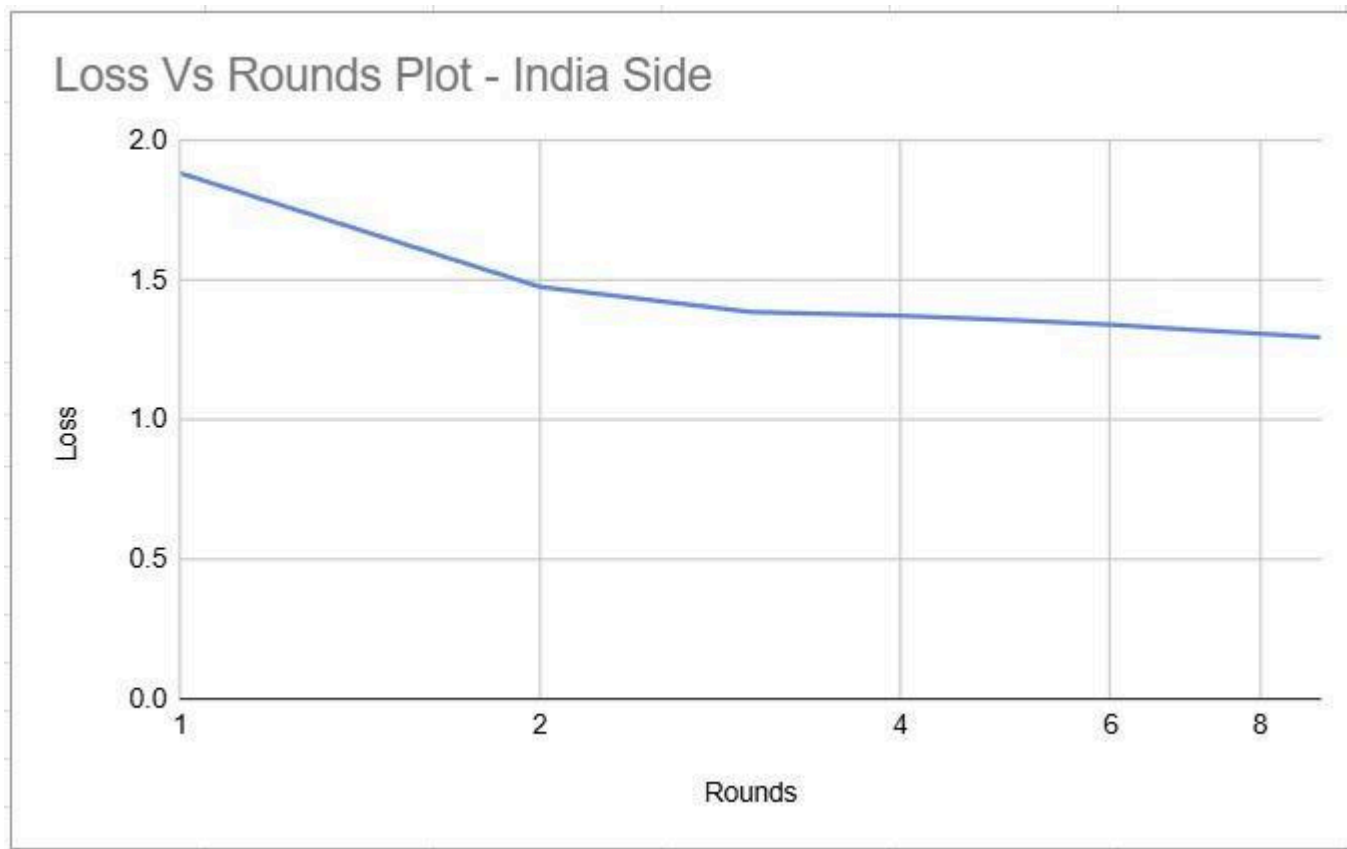
Item	Setting
Federation ID	slakshna-M0
Expected participants	2
Federated clock	300 seconds
Synchronization deadline	300 seconds
Trust update coefficient	0.1
Public-network path	Explicit authenticated peer through a Playit UDP tunnel
Delta sparsity	Top 10%
Quantization	Symmetric INT8
Typical Base64 model-update payload	Approximately 5.43 MiB
Raw training data transferred	None
Full model weights transferred	None

Training trajectory

The Indian endpoint joined the federation at 16:30 IST , and the first training boundary began at 16:55.

We completed 9 rounds of federated training successfully, and found the loss to be converging.



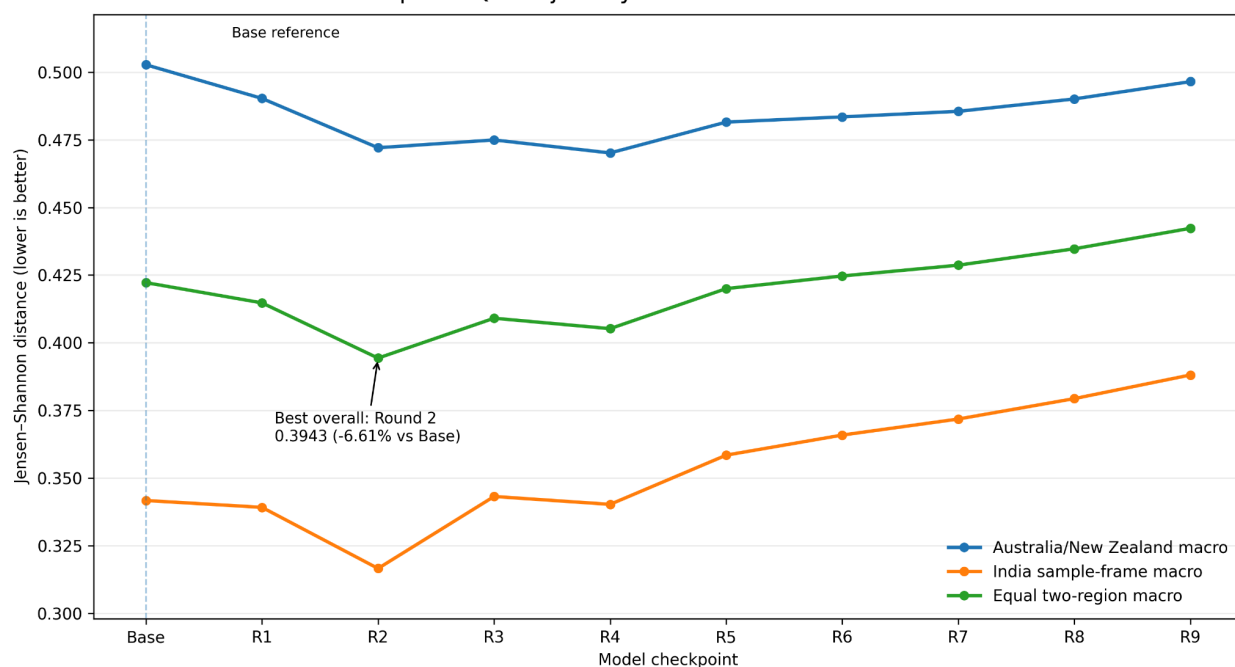


GlobalOpinionQA evaluation

Evaluation used the standalone shared GOQA package and loaded the base model plus each Round 1–9 LoRA adapter directly, without merging adapters into full model copies. The evaluation retained 1,106 questions with at least one Australia, New Zealand, or Indian human distribution. Five deterministic option orders were used for each question, producing 5,530 prompts per model state and 1,831 valid country/sample-frame comparison pairs.

Jensen–Shannon distance compares the model's option distribution with the available human response distribution; lower values are better. The Australia/New Zealand metric is the equal macro average of the two country means. The India metric is the equal macro average of the current-national, non-national, and old-national sample frames. The primary two-region metric weights the Australia/New Zealand and India values equally. Missing human groups are omitted for the corresponding question rather than imputed.

GlobalOpinionQA Trajectory — Base + Federated Rounds 1–9



Model state	Australia/NZ JSD	India JSD	Two-region JSD	Australia/NZ Δ vs Base (%)	India Δ vs Base (%)	Two-region Δ vs Base (%)
Base	0.503	0.342	0.422	0.000	0.000	0.000
Round 1	0.490	0.339	0.415	-2.476	-0.740	-1.773
Round 2	0.472	0.317	0.394	-6.103	-7.353	-6.609
Round 3	0.475	0.343	0.409	-5.528	0.436	-3.115
Round 4	0.470	0.340	0.405	-6.484	-0.417	-4.029
Round 5	0.482	0.358	0.420	-4.216	4.907	-0.525
Round 6	0.484	0.366	0.425	-3.833	7.080	0.583
Round 7	0.486	0.372	0.429	-3.425	8.815	1.527
Round 8	0.490	0.379	0.435	-2.515	11.025	2.963
Round 9	0.497	0.388	0.442	-1.239	13.581	4.757

Round 2 is the best checkpoint on the equal two-region metric and on the India component. Round 4 is narrowly best on the Australia/NZ component.

Observations

The run terminated before reaching the planned 20 rounds because the sites became mutually offline. No Slakshna or Bhaskera crash was observed in association with the failure, and its precise cause remains undetermined. The Playit service or the surrounding public-network communication path is therefore identified as the primary area for further investigation.

Hyperparameters were independently validated on both sides without issue. For future federated learning experiments, synchronization of round execution times across nodes should be treated as an essential component of the experimental setup. Furthermore, deterministic behaviour requires accounting for the storage overhead associated with retaining all round states. An alternative approach would be to employ a more storage-efficient algorithm for merging deltas; however, such an approach may introduce non-deterministic behaviour. The trade-off between storage efficiency and determinism should therefore be systematically evaluated in subsequent experiments.

Round 9 constitutes the latest valid federated artifact and should be retained for joint-run comparisons. Round 2 is the preferred checkpoint for GOQA from this trajectory, with a two-region JSD of 0.394, corresponding to a 6.609% reduction relative to base. Although a longer rerun remains necessary to obtain a stable final model, the present run provides sufficient evidence that the cross-country training and evaluation pipeline is operational and capable of completing the intended experimental workflow.

Conclusion

The session successfully met the practical M0 objective: two geographically separate sites trained the agreed model on their respective private regional data, exchanged compressed LoRA deltas, retained valid synchronized adapters, and demonstrated measurable improvements on the shared GOQA evaluation. The training stack remained stable throughout the session.