

Model parallelism with Pytorch

Gyula Ujlaki

DKF Kft., NCC Hungary – HPC Scientific expert

Use cases for model parallelism

Calculating model memory requirements:

Inference:

$$\text{num_para} * \text{precision} / 8\text{bit} = \text{mem_req}$$

$$\text{mem_req} * 1.2 = \text{inference_mem_req}$$

Training:

$$\text{optim_states}_{\text{AdamW}} = 4 \times \text{mem_req}$$

$$\text{optim_states}_{\text{SGD}} = 1.5 \times \text{mem_req}$$

$$\text{gradients} = \text{mem_req}$$

$$\text{activations} = \text{sbhL}$$

$$\text{activations} = \text{sbhL}(10 + 24/t)$$

$$\text{activations} = \text{sbhL}(10 + 24/t + (5a*s)/h*t)$$

s – sequence length in tokens

b – batch size per GPU

h – dimension of the hidden size within each transformer layer

L – the number of layers in the transformer model

a – the number of attention heads in the transformer model

t – degree of tensor parallelism being used (1 if not)

No sequence parallelism being used

No activations stored in fp16

Training memory requirements of Llama 70B + AdamW optimizer, without calculating sequence lengths, batch size, etc.:

~70 * 10⁹ parameters, fp16 precision

70 * 10⁹ * 2 (bytes) = 140 * 10⁹ (model memory requirements)

140 * 10⁹ * 4 (bytes) = 560 * 10⁹ (with AdamW optimizer states)

560 * 10⁹ + 140 * 10⁹ = 700 * 10⁹ bytes (with gradients)

Activation size depends on a lot of factors, thus we will not approximate this size. Even without that:

700 / 1024³ = 652 GB (memory requirements) * 1.2 (for overhead)

Officially 500 GB: <https://huggingface.co/blog/llama31>

Depends on specific implementations

Source: <https://blog.eleuther.ai/transformer-math/>

Pytorch - Layer-wise model parallelism

- a. If only 1 GPU is available: CPU offloading, a few layers run on CPU
- b. If at least 2 GPU available: divide layers evenly between the two GPUs.

Only 1 model, running on multiple devices, increasing/multiplying the memory available for the model!
Slower than 1 GPU, because of the communication overhead between GPUs.

Model can be distributed even among nodes as well, with Remote Procedure Calls: <https://docs.pytorch.org/docs/stable/rpc.html>

Cases when implementing an LWMP strategy is useful:

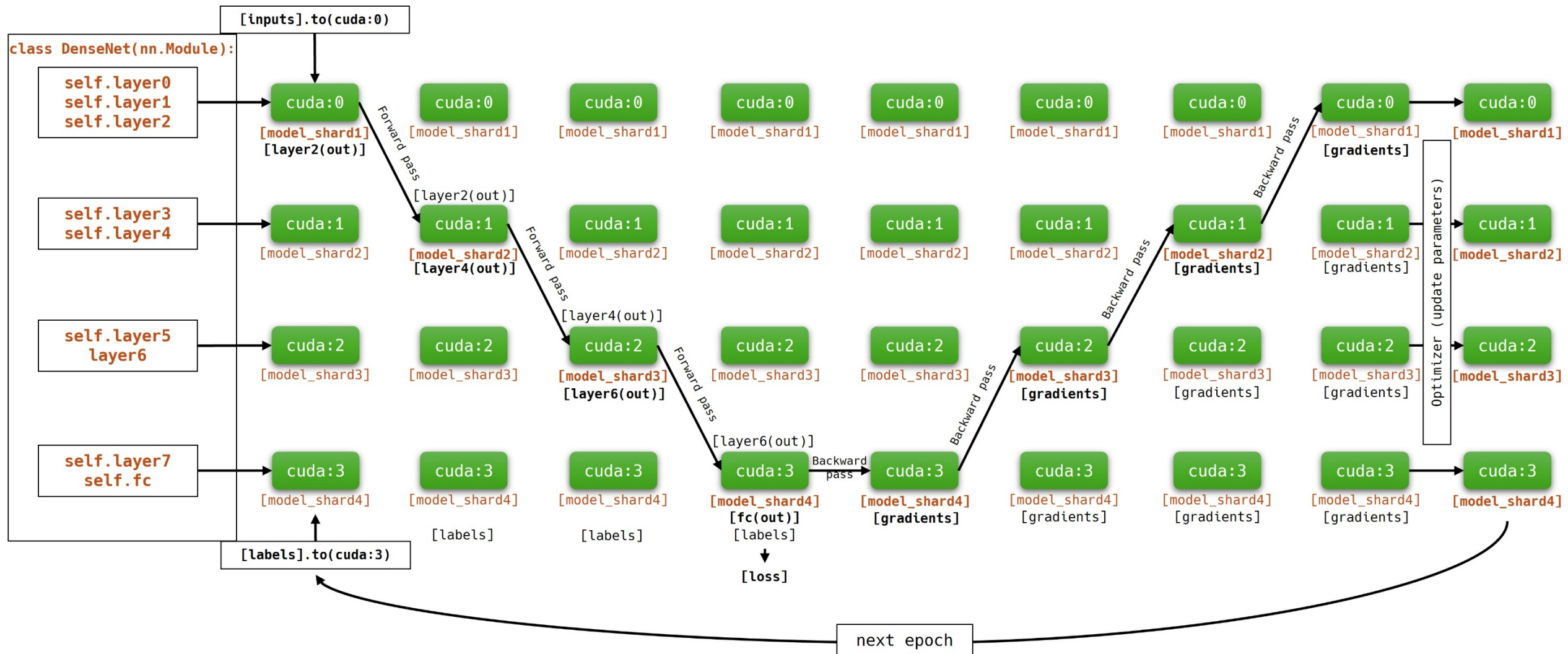
- c. Model size is slightly larger than GPU memory.
- d. If pipeline parallel logic implemented – no exercise for this part, as this is significantly more difficult to implement, than LWMP
- e. Large language models cannot fit into a single GPU, and e.g. Tensor Parallelism also used for in-layer distribution.

LWMP can be combined with other parallelism techniques as well:

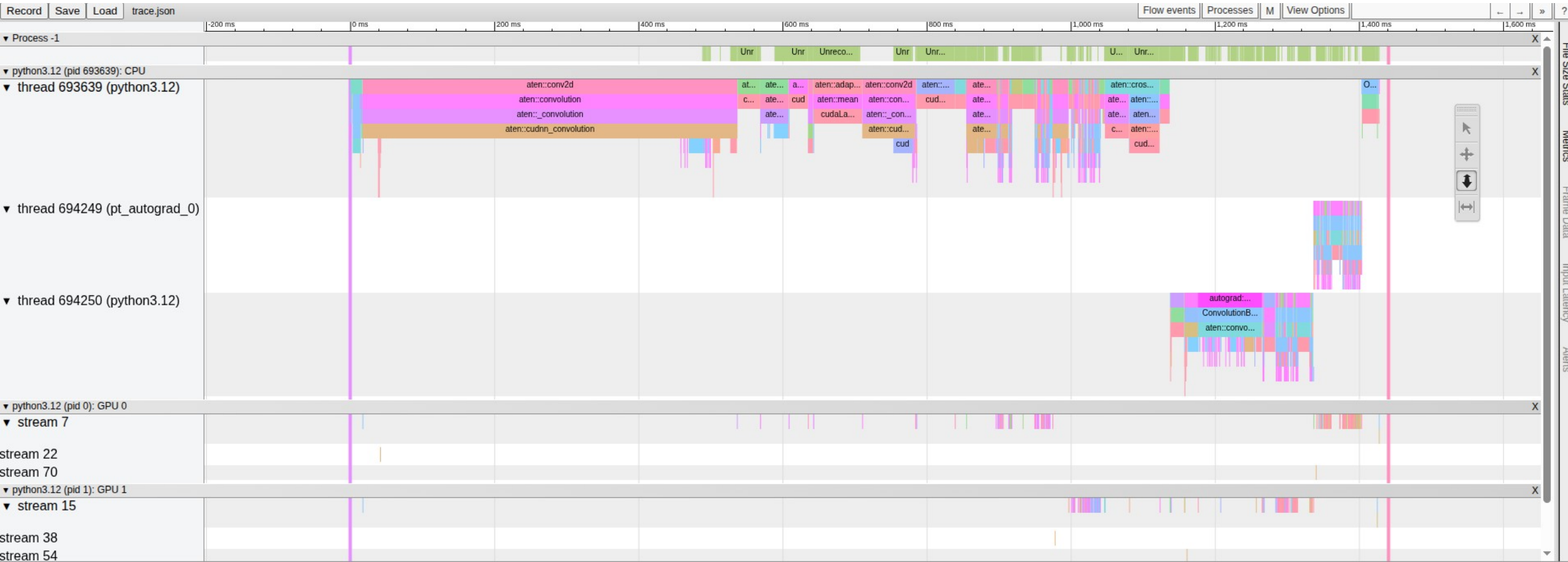
- Distributed Data Parallel – DDP → inefficiency upscaled
- Pipeline parallelism

etc.

Layer-wise model parallelism



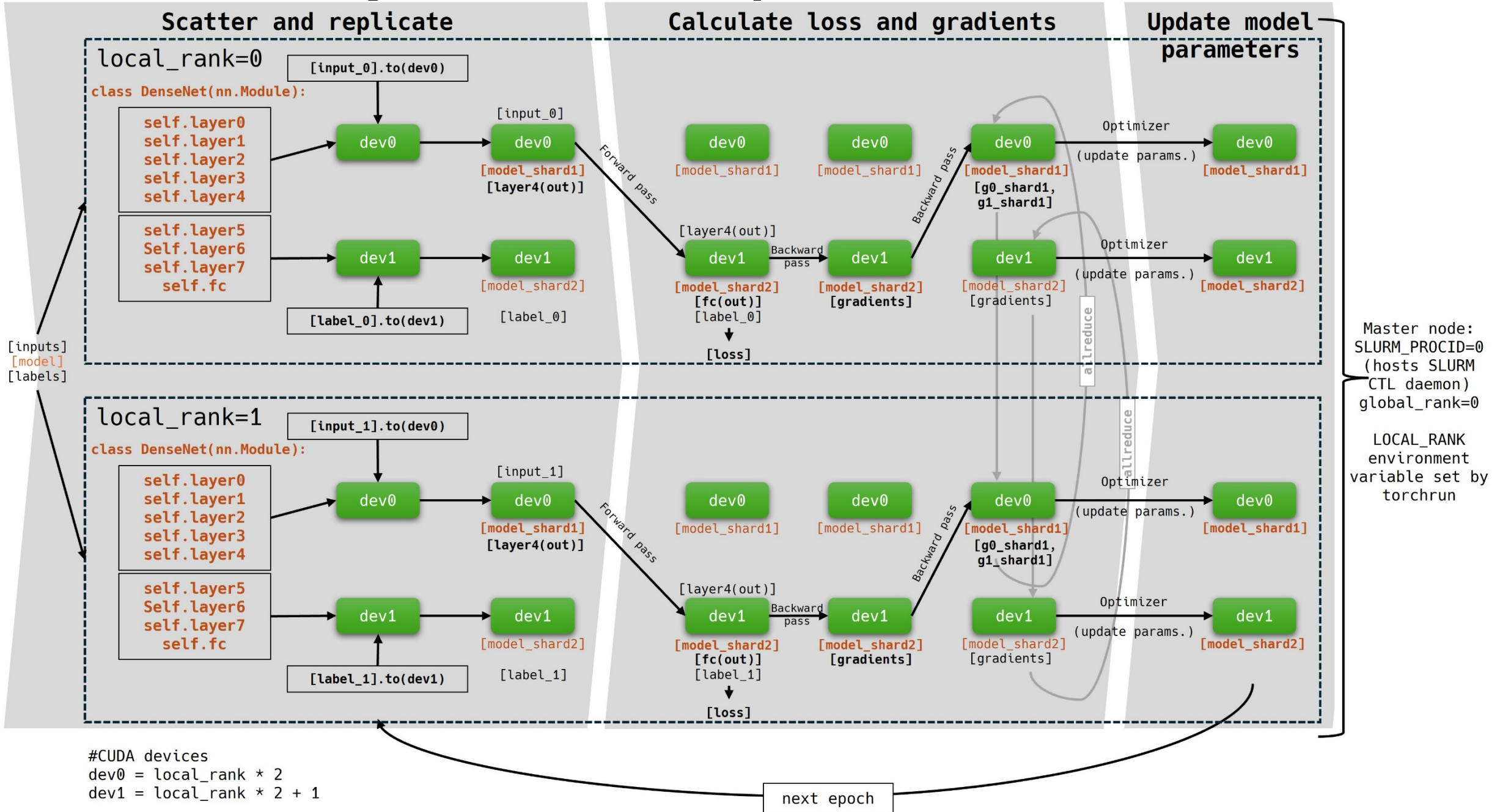
Tracing operations and tensors during training



1 item selected.		Slice (1)
Title	autograd::engine::evaluate_function: ConvolutionBackward0 🔍	
Category	cpu_op	
User Friendly Category	other	
Start	1,176.511 ms	
Wall Duration	89.303 ms	
Self Time	0.017 ms	
▼Args		
External id	4165	
Record function id	0	
Sequence number	291	
Fwd thread id	1	
Ev Idx	68	

ssh -L 16001:localhost:16001 youruser@login.Leonardo.cineca.it
tensorboard --logdir=/home/user/log/path --port 16001 --host localhost
(everyone has to use a unique port number)
Browser: localhost:16001

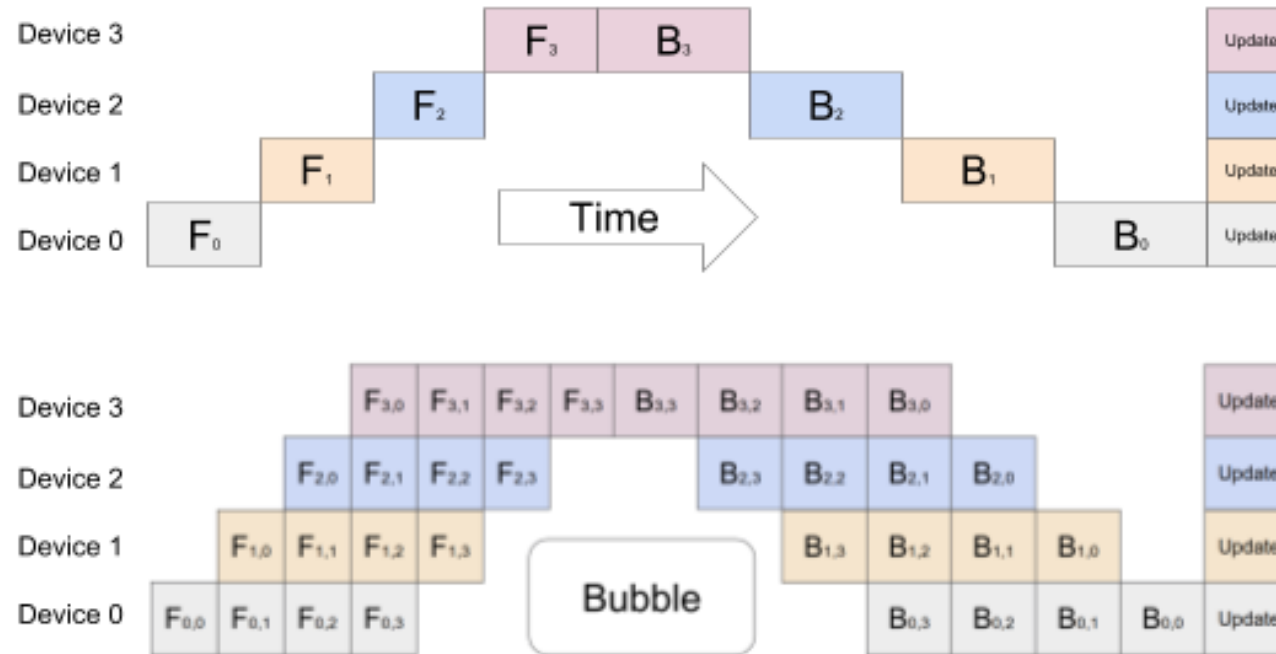
Layer-wise model parallelism + DDP



Basic model parallel techniques with Pytorch

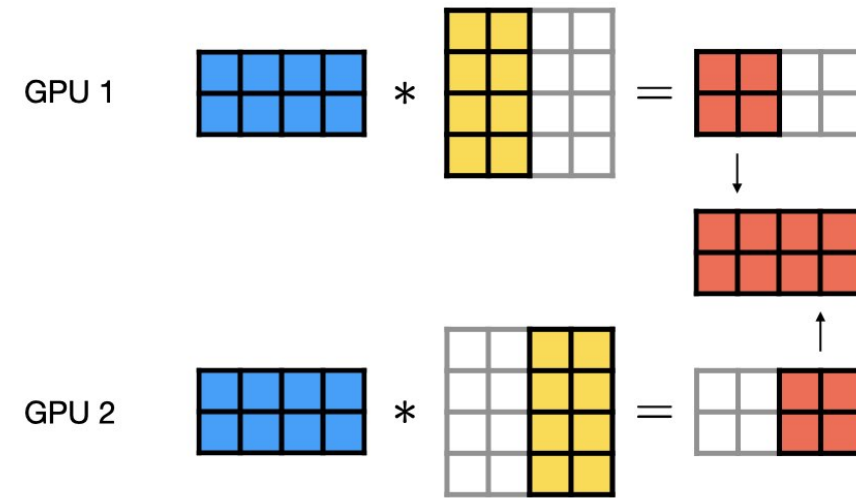
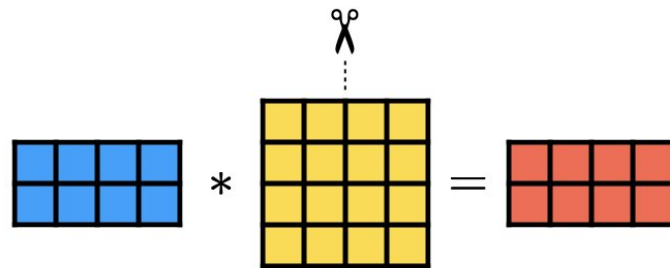
Large amount of data	Too large model for 1 GPU	Both
No parallelism, infinite time	Layer-wise model parallel, or naïve model parallel	FSDP – Fully Sharded Data Parallel
DP – Data Parallel	TP - Tensor parallel	DDP + PP or TP
DDP – Distributed Data Parallel	PP - Pipeline parallel	+ any other combination of methods

Pipeline Parallelism



Top: The naive model parallelism strategy leads to severe underutilization due to the sequential nature of the network. Only one accelerator is active at a time. Bottom: GPipe divides the input mini-batch into smaller micro-batches, enabling different accelerators to work on separate micro-batches at the same time.

Tensor Parallelism in principle



Thank you!