

Accelerated training through iterative gradient propagation along the residual path

Erwan Fagnou¹, Paul Caillon¹, Blaise Delattre^{1,2}, Alexandre Allauzen^{1,3}

¹ Miles Team, LAMSADE, Université Paris Dauphine-PSL, Paris, France

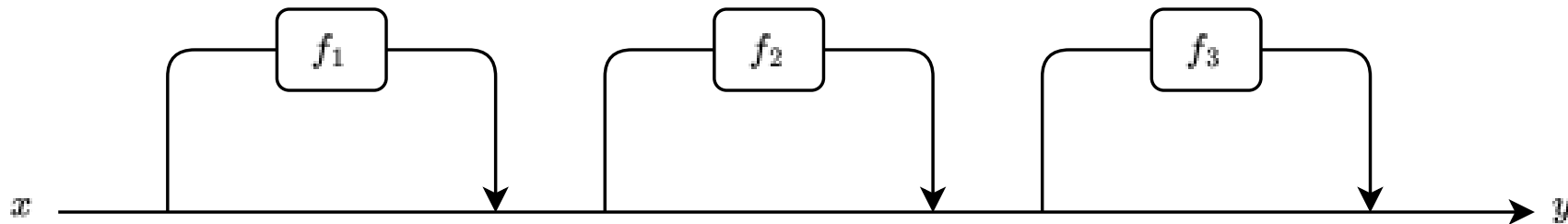
² Foxstream, Vaulx-en-Velin, France

³ ESPCI PSL, Paris, France

Residual models

- Simplest form: identity residual connection

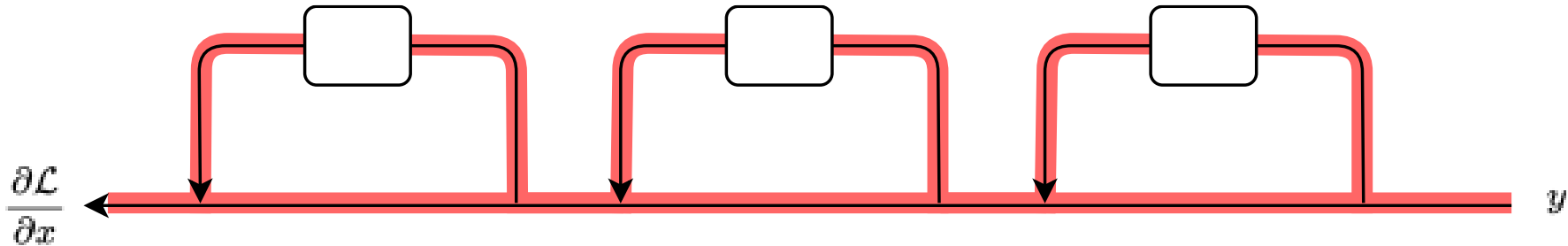
$$\begin{aligned} h_i &= h_{i-1} + f_i(h_{i-1}) \\ &= x + \sum_{j=1}^i f_j(h_{j-1}) \end{aligned}$$



Residual models

- Backpropagation:

$$\begin{aligned}\frac{\partial \mathcal{L}}{\partial h_i} &= \frac{\partial \mathcal{L}}{\partial h_{i+1}} (I + J_{j+1}) \\ &= \frac{\partial \mathcal{L}}{\partial h_L} (I + J_L) \dots (I + J_{i+1})\end{aligned}$$



Gradient decomposition

$$\frac{\partial \mathcal{L}}{\partial h_i} = \frac{\partial \mathcal{L}}{\partial h_L} (I + J_L) \dots (I + J_{i+1})$$

Gradient decomposition

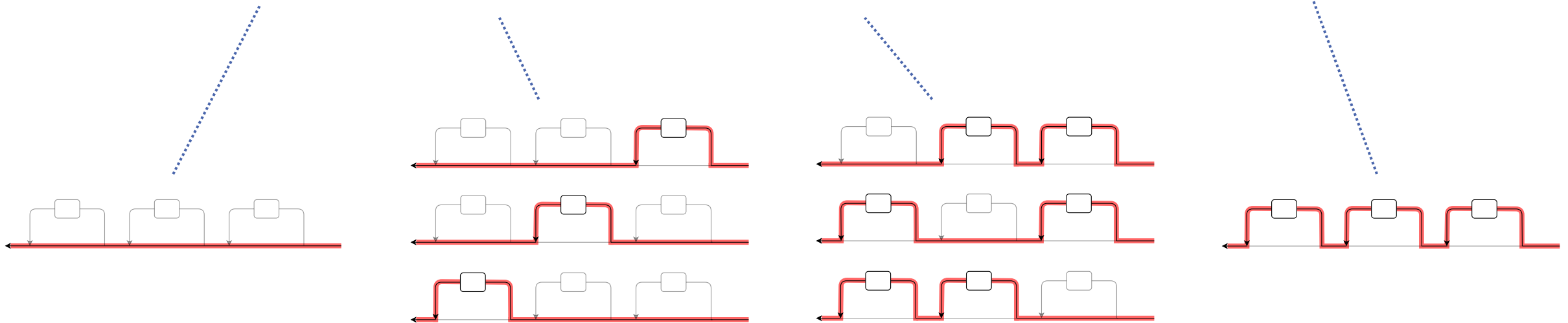
$$\frac{\partial \mathcal{L}}{\partial h_i} = \frac{\partial \mathcal{L}}{\partial h_L} (I + J_L) \dots (I + J_{i+1})$$

$$= \frac{\partial \mathcal{L}}{\partial h_L} \left[I + \sum_{j=i+1}^L J_j + \sum_{j=i+1}^L \sum_{k=j+1}^L J_j J_k + \dots + J_L J_{L-1} \dots J_{i+1} \right]$$

Gradient decomposition

$$\frac{\partial \mathcal{L}}{\partial h_i} = \frac{\partial \mathcal{L}}{\partial h_L} (I + J_L) \dots (I + J_{i+1})$$

$$= \frac{\partial \mathcal{L}}{\partial h_L} \left[I + \sum_{j=i+1}^L J_j + \sum_{j=i+1}^L \sum_{k=j+1}^L J_k J_j + \dots + J_L J_{L-1} \dots J_{i+1} \right]$$



Gradient decomposition

True gradient:

$$\frac{\partial \mathcal{L}}{\partial h_i} = \sum_{n=0}^L \text{"gradients going through } n \text{ blocks"}$$

Approximation with partial sum:

$$w_i^k = \sum_{n=0}^k \text{"gradients going through } n \text{ blocks"}$$

Gradient decomposition

True gradient:

$$\frac{\partial \mathcal{L}}{\partial h_i} = \sum_{n=0}^L \text{"gradients going through } n \text{ blocks"}$$

Approximation with partial sum:

$$w_i^k = \sum_{n=0}^k \text{"gradients going through } n \text{ blocks"}$$

→ Why would this work?

- Gradient vanishing reduces the gradients going through many blocks

Highway backpropagation

Iterative computation:

$$\begin{cases} w_i^0 &= \sum_{j=i}^L \frac{\partial \mathcal{L}_j}{\partial h_j} \\ w_i^{k+1} &= w_i^0 + \sum_{j=i+1}^L w_j^k J_j \end{cases}$$

Highway backpropagation

Iterative computation:

$$\begin{cases} w_i^0 = \sum_{j=i}^L \frac{\partial \mathcal{L}_j}{\partial h_j} \\ w_i^{k+1} = w_i^0 + \sum_{j=i+1}^L w_j^k J_j \end{cases}$$

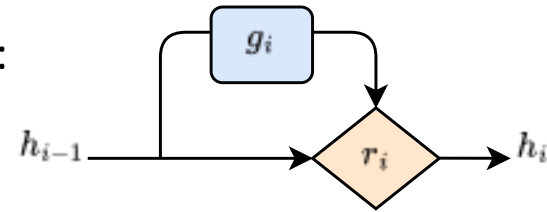
1) Parallel backpropagation through each block

2) Cumulative sum (parallelizable)

General framework

In the paper, we extend the framework to a broader set of models:

- Arbitrary residual connection functions:

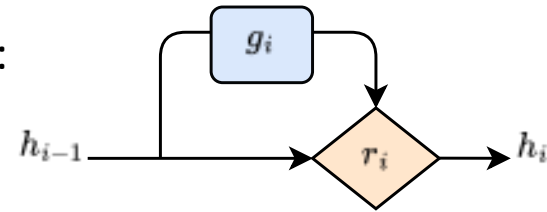


→ ResNets, post-layernorm transformers, ...

General framework

In the paper, we extend the framework to a broader set of models:

- Arbitrary residual connection functions:



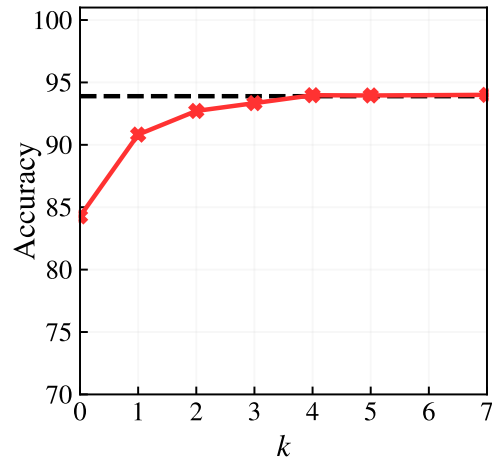
→ ResNets, post-layernorm transformers, ...

- Handle intermediary losses:

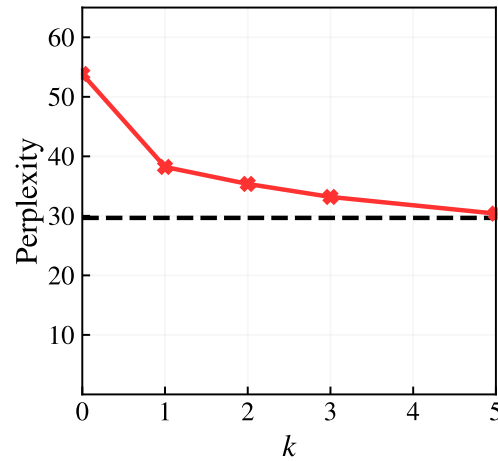
$$\mathcal{L}(x) = \sum_{i=0}^L \mathcal{L}_i(h_i)$$

→ RNNs: GRUs, LSTMs, ...

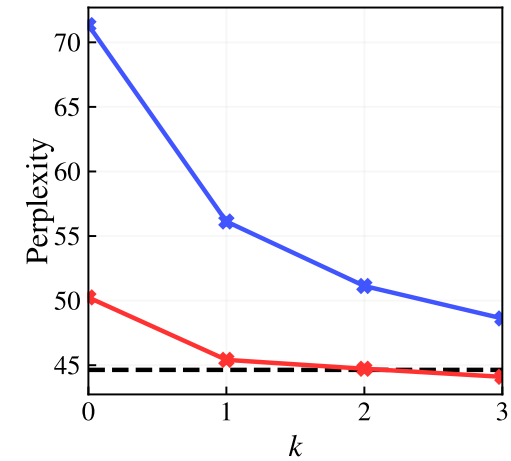
Result



ResNet110
CIFAR10
($L = 54$)



Transformer
Wikitext103
($L = 24$)



RNN (GRU x3)
Wikitext103
($L = 256$)

→ Matches backpropagation with $k \ll L$ iterations.

Thank you!

Contact: erwan.fagnou@dauphine.psl.eu

Paper:

