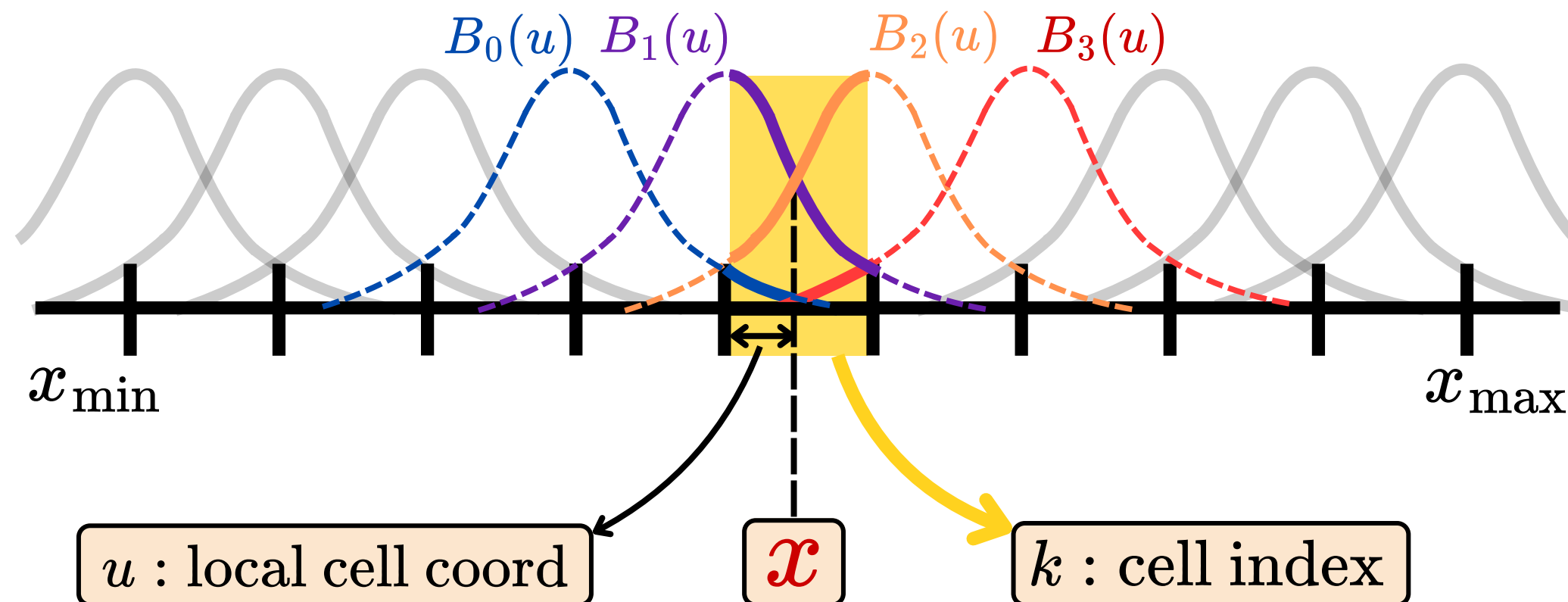


Ultrafast On-Chip Online Learning via Spline Locality in Kolmogorov-Arnold Networks

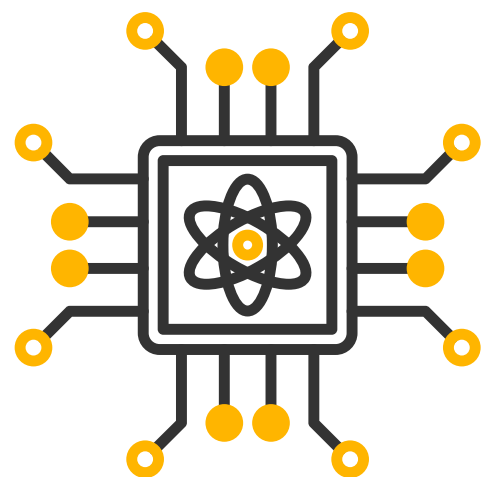


Duc Hoang*, Aarush Gupta*, Philip Harris

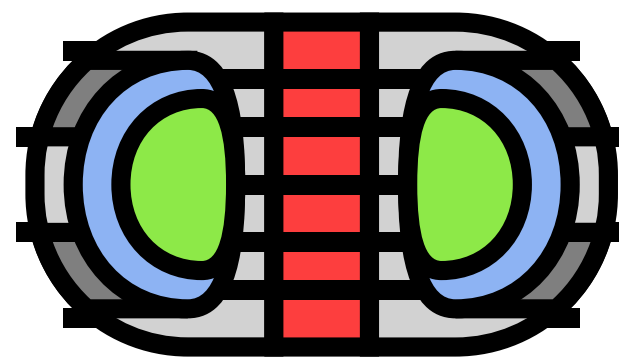
MIT

*equal contribution





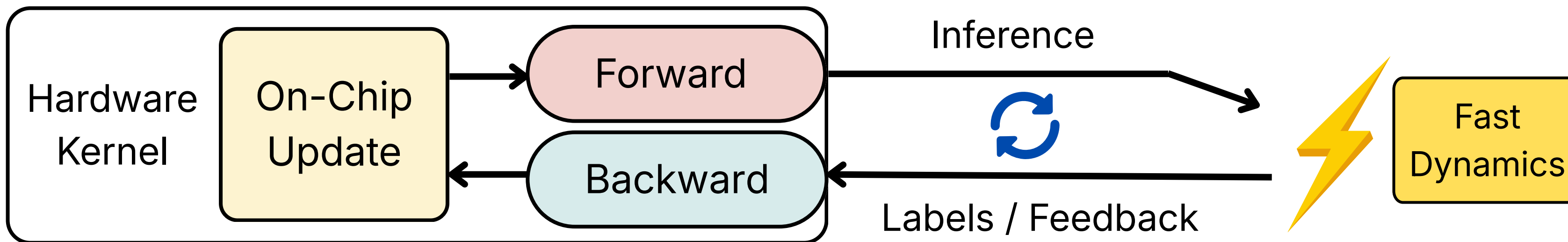
Quantum Computing

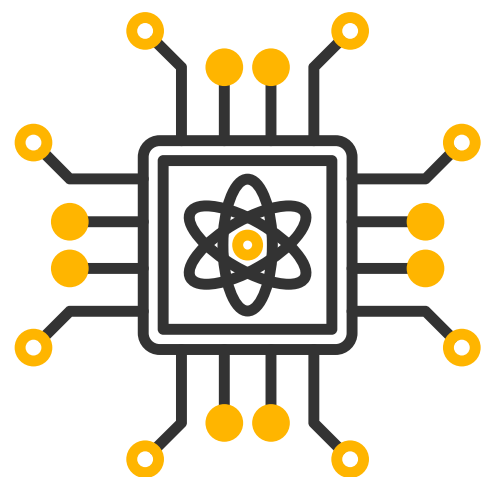


Fusion Control

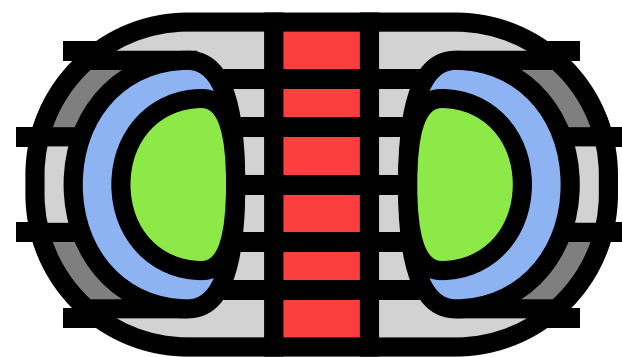


Communication





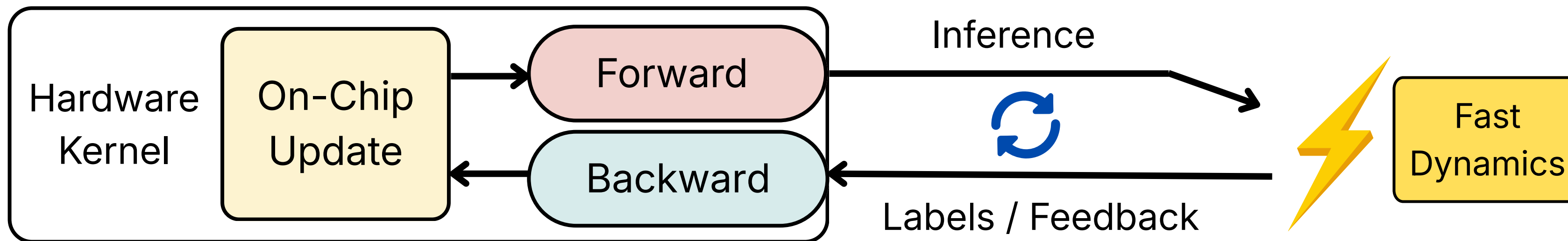
Quantum Computing



Fusion Control

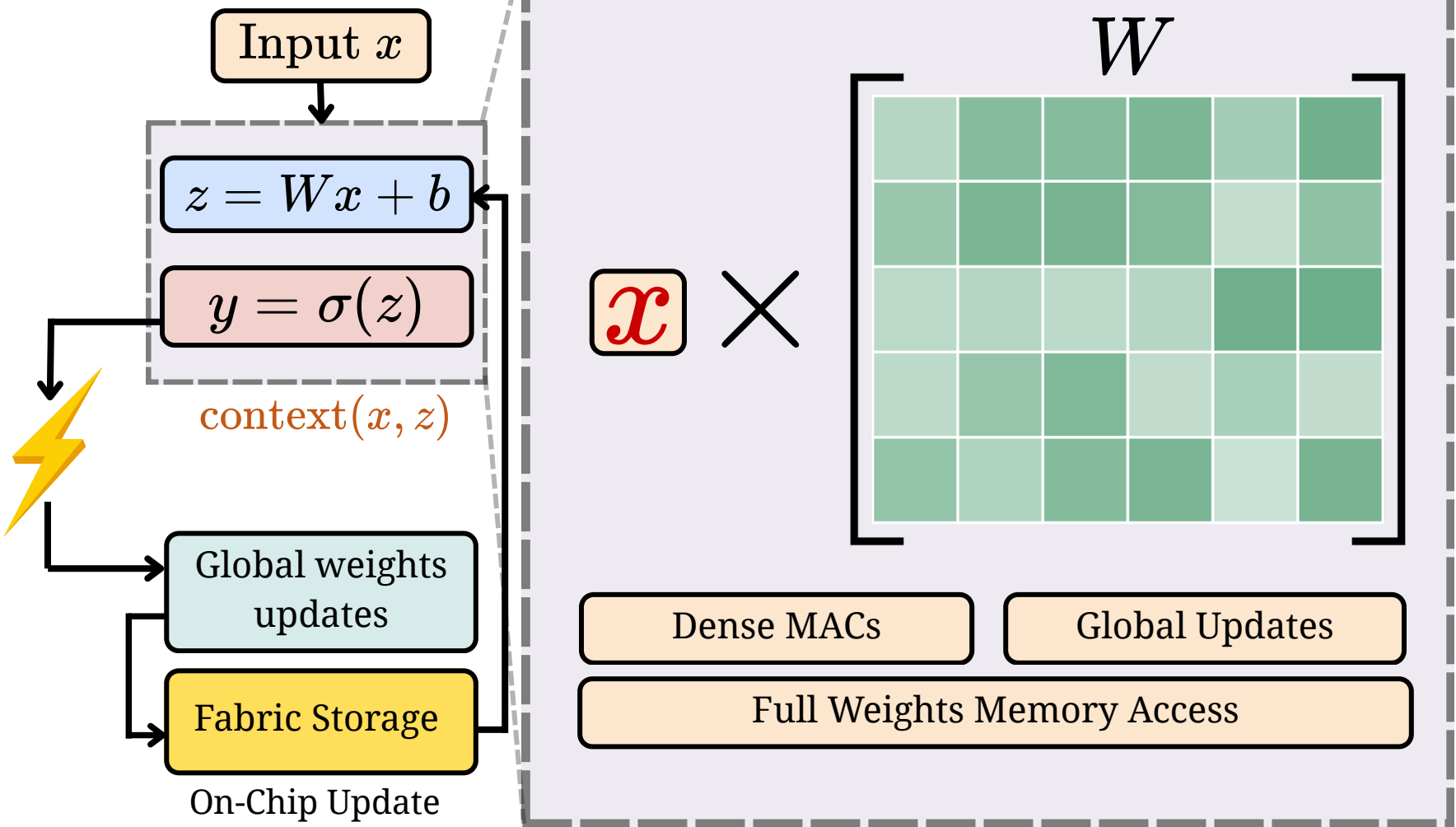


Communication



Sub-microsecond update latency needed! (FPGA / ASIC)

MLP Implementation



MLP $y_q = \phi \left(\sum_p w_{q,p} x_p \right)$

MLP $y_q = \phi \left(\sum_p w_{q,p} x_p \right)$

learned activations replace learned weights

KAN $y_q = \sum_p \phi_{q,p}(x_p)$

MLP $y_q = \phi \left(\sum_p w_{q,p} x_p \right)$

learned activations replace learned weights

KAN $y_q = \sum_p \phi_{q,p}(x_p)$

where $\phi_{q,p}(x) = \sum_i c_{q,p,i} B_i(x)$

we learn these coefficients

MLP $y_q = \phi \left(\sum_p w_{q,p} x_p \right)$

learned activations replace learned weights

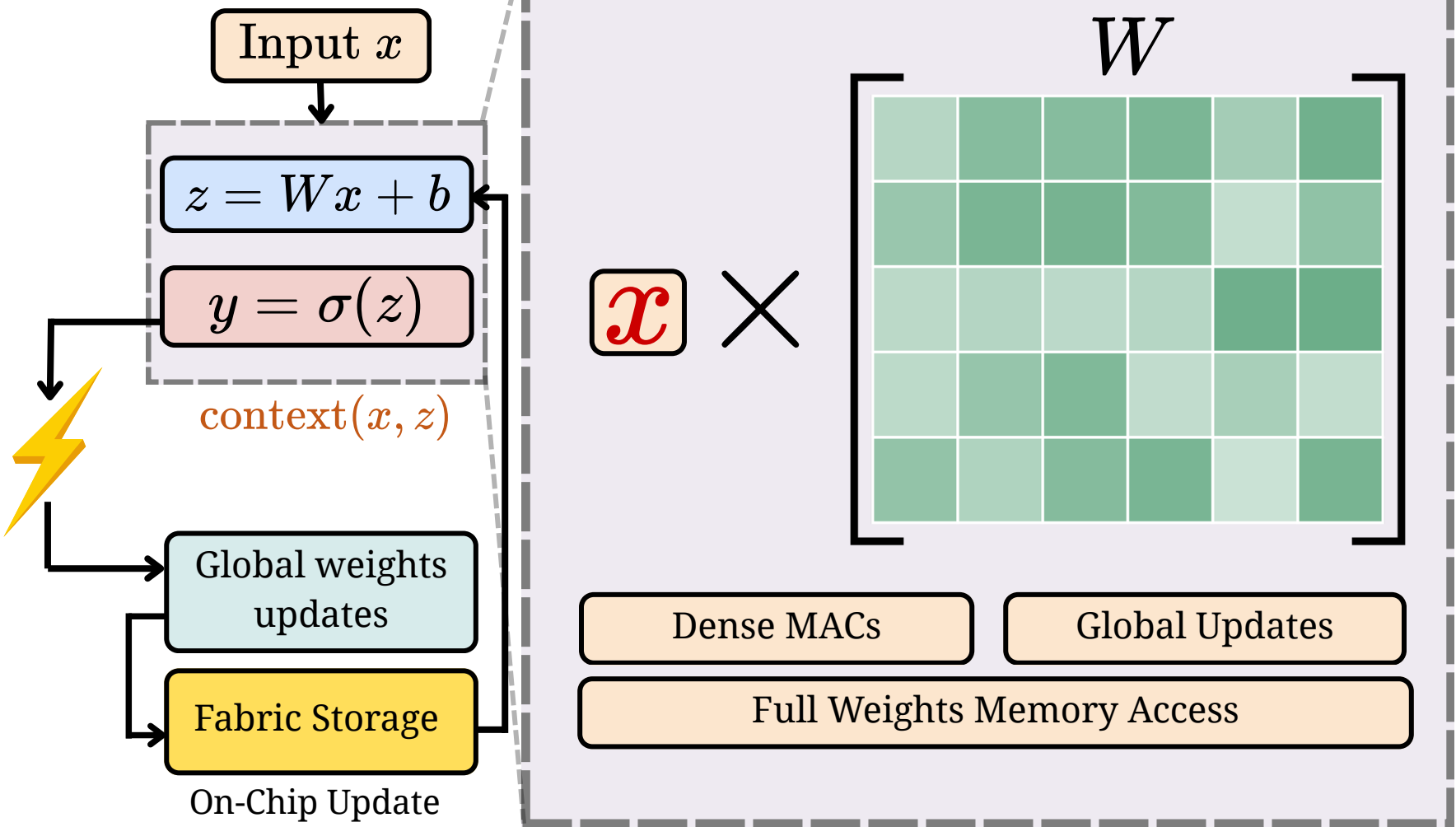
KAN $y_q = \sum_p \phi_{q,p}(x_p)$

where $\phi_{q,p}(x) = \sum_i c_{q,p,i} B_i(x)$

B-splines

- $\{B_i\}$ are **local** → small fraction computed per input
- $\{B_i\}$ are bounded → well-behaved under quantization

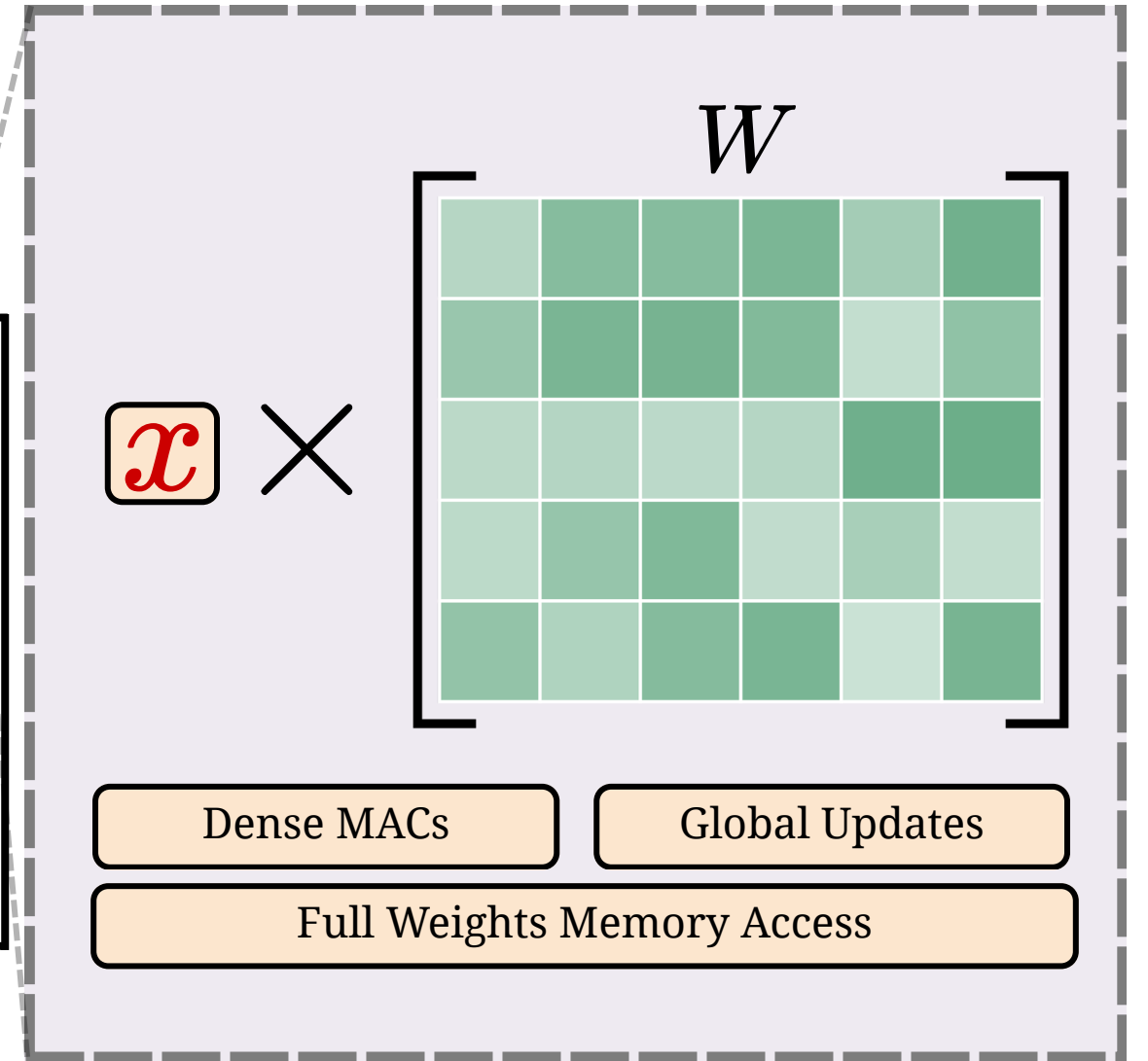
MLP Implementation



Spline order $p=3$

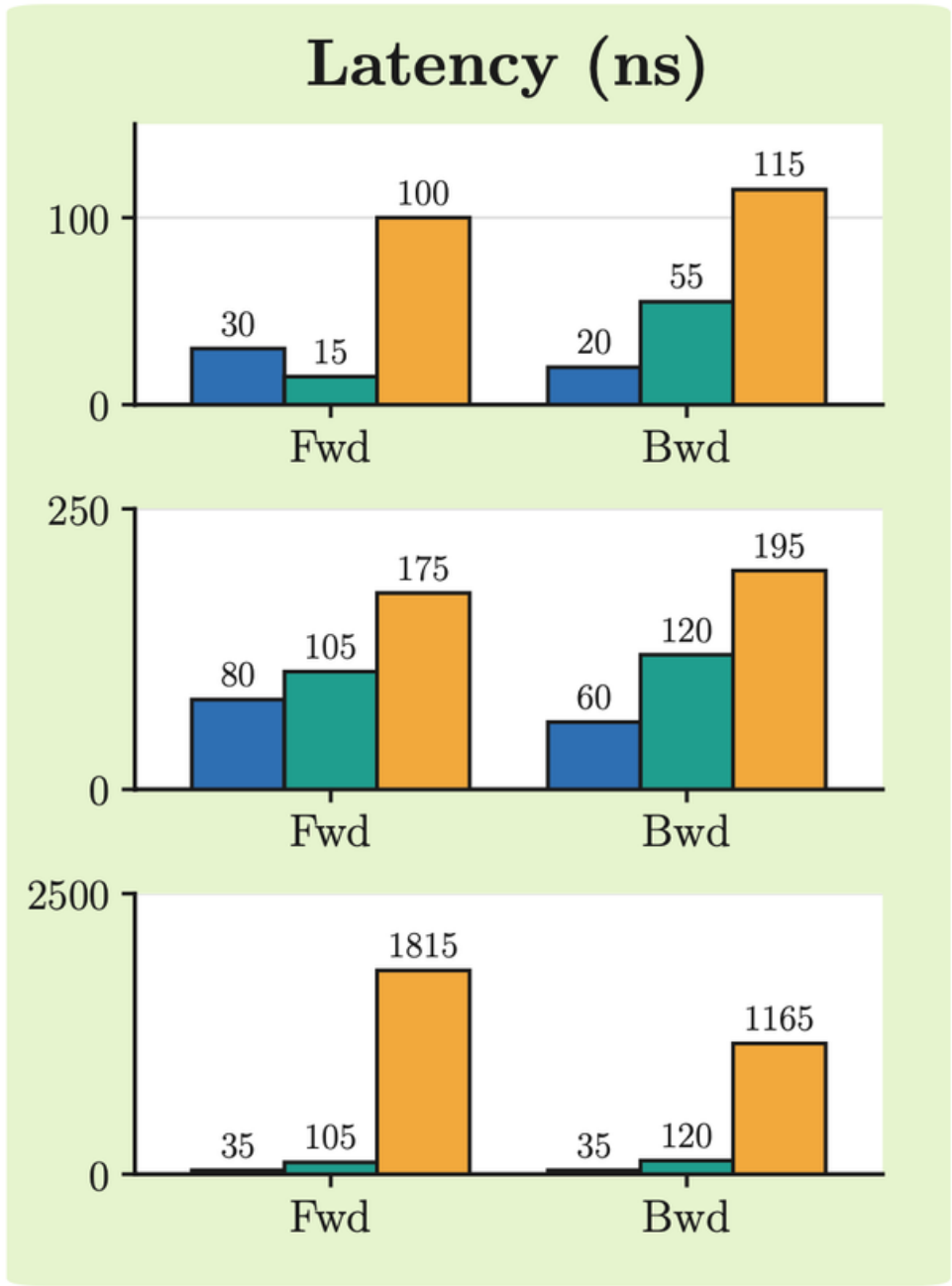
$$\phi(x) = \sum_{r=0}^{p=3} \text{Sparse active coeffs } w_r \text{ Pre-computed LUT } B_r(u)$$

The diagram illustrates the evaluation of a cubic B-spline function $\phi(x)$ at a point x . The domain is divided into cells by a grid. The basis functions $B_0(u)$, $B_1(u)$, $B_2(u)$, and $B_3(u)$ are shown as dashed curves. A yellow box highlights the active region around the point x , where the function value is computed as a weighted sum of the active basis functions. The point x is mapped to a local cell coordinate u and a cell index k .

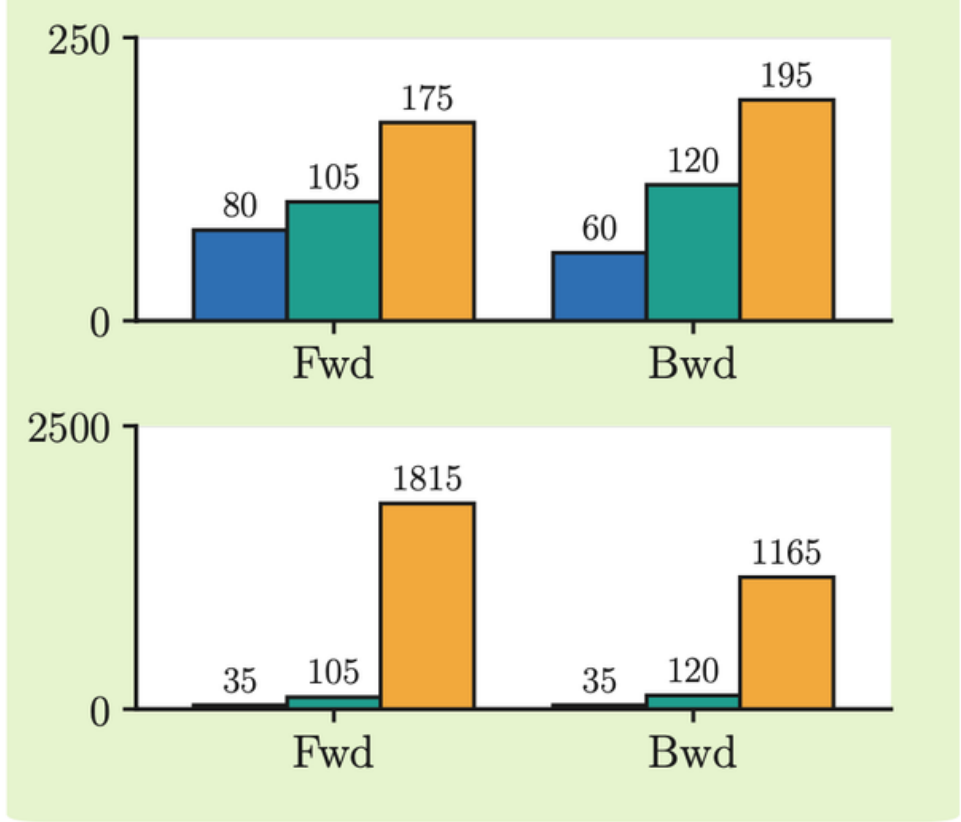




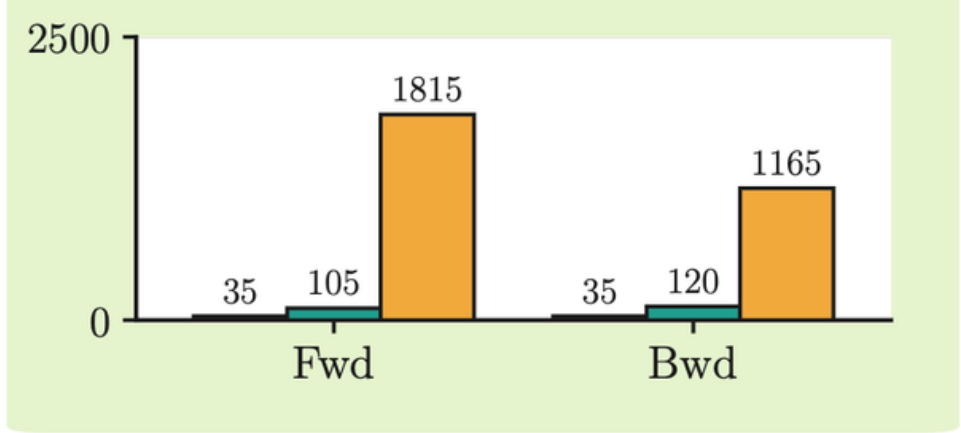
Adaptive
Func. Approx.



Single-Shot
Qubit Readout

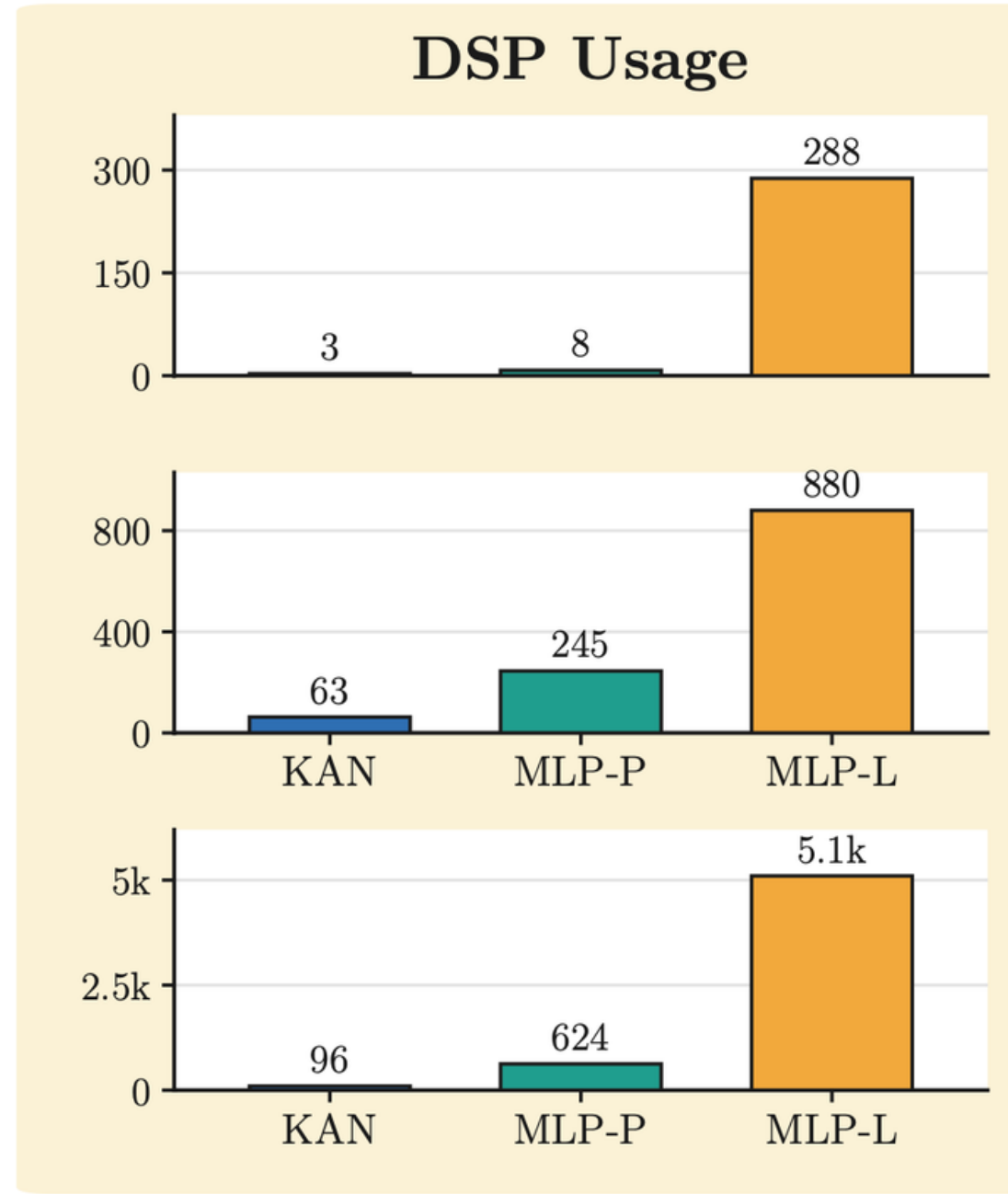


Acrobot

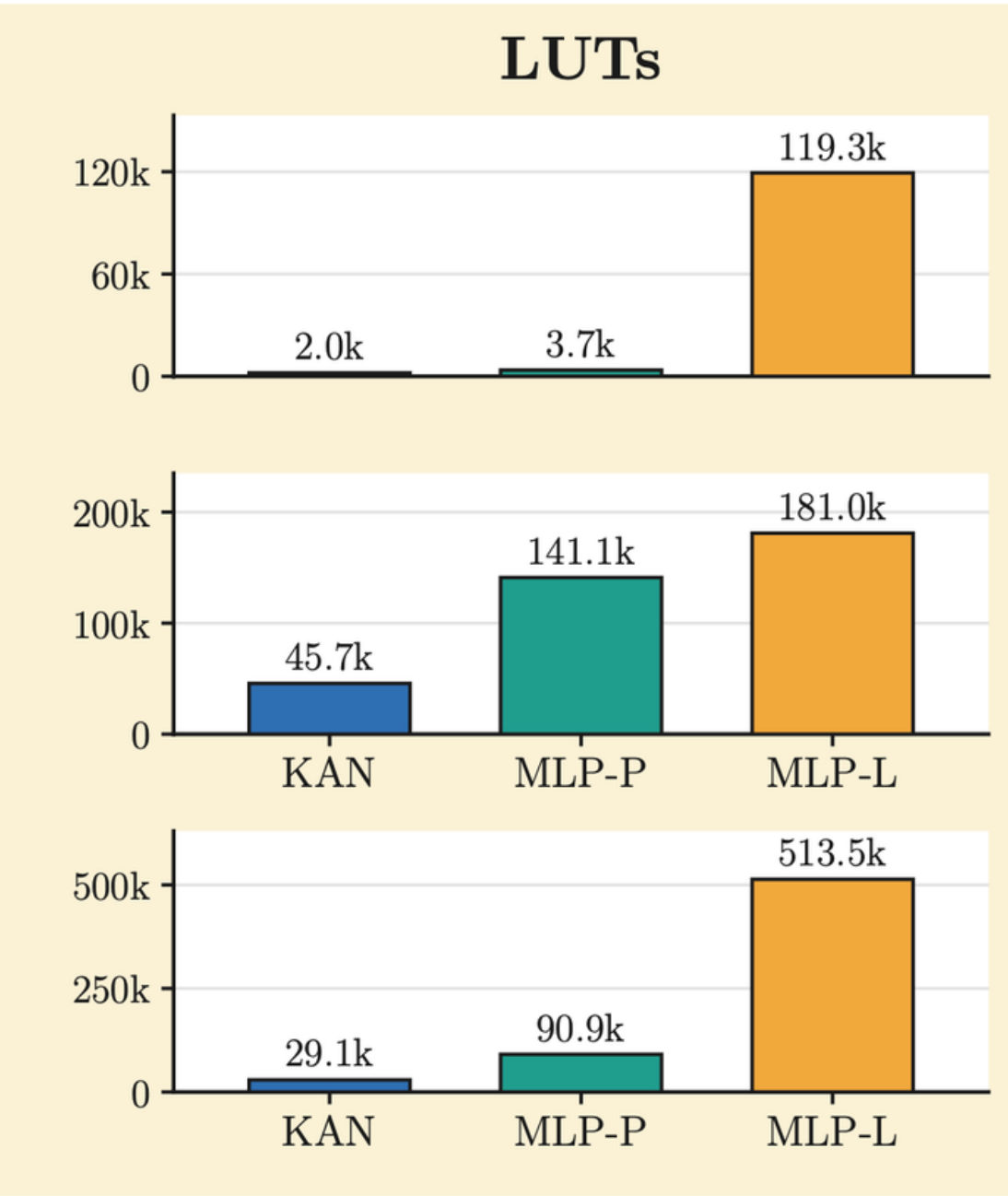


Much faster latency

DSP Usage



LUTs



Superior on-chip resource scaling

KAN learners scale to 50k+ params with sub-microsecond latencies!