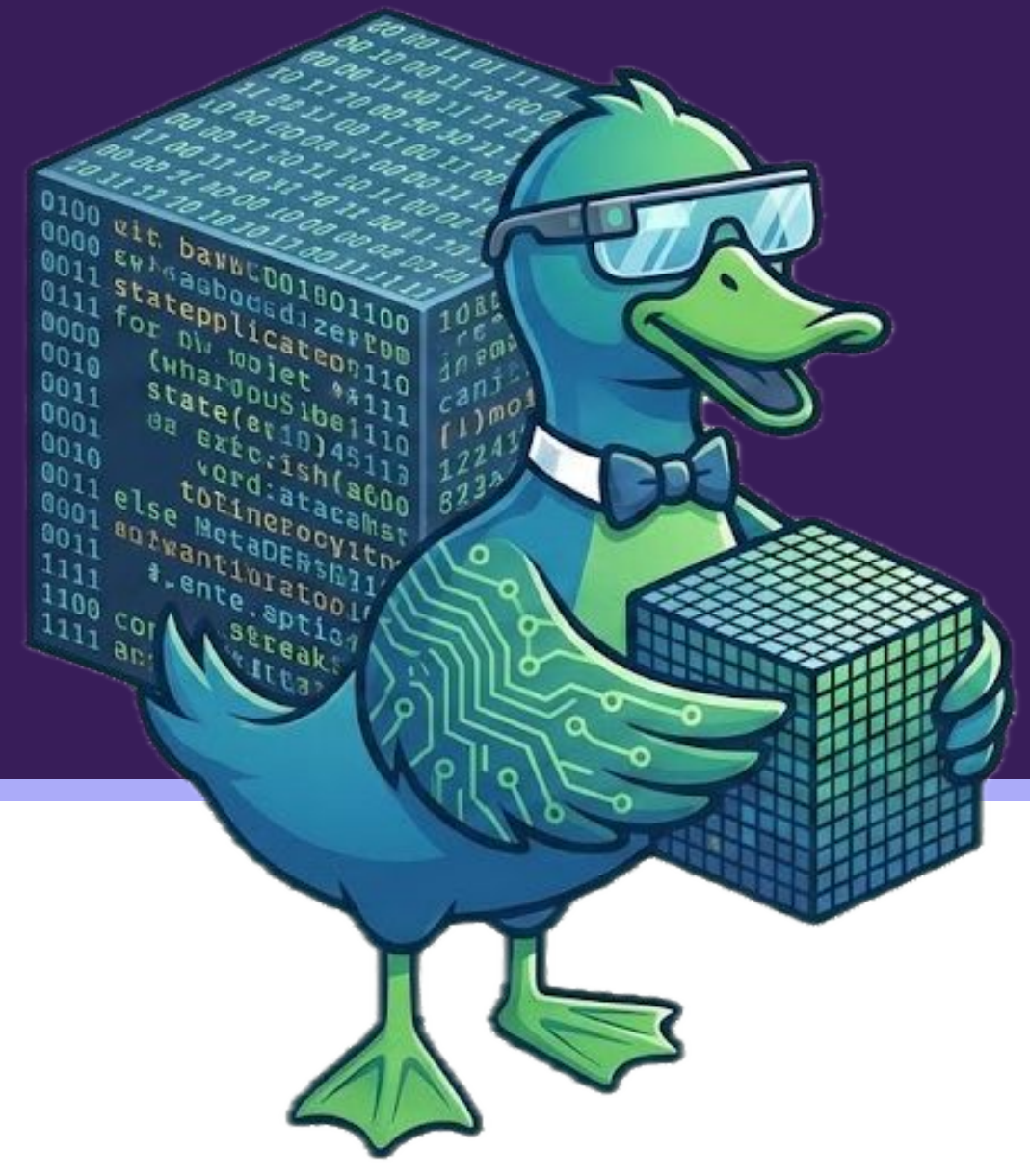


Float8@2bits: Entropy Coding Enables Data-Free Model Compression

Patrick Putzky*, Martin Genzel*, Mattes Mollenhauer, Sebastian Schulze, Thomas Wollmann, Stefan Dietzel

Merantix Momentum GmbH, Berlin, Germany



We introduce EntQuant, a novel compression method that keeps model weights in 8-bit precision but optimizes their distribution to enjoy low entropy. Using GPU-accelerated entropy codecs, we obtain compressed models that remain functional up to effectively 2 bits, without calibration or fine-tuning.

< 10 min
to compress a 70B model

Compress as you
deploy the model.
No GPU Cluster.

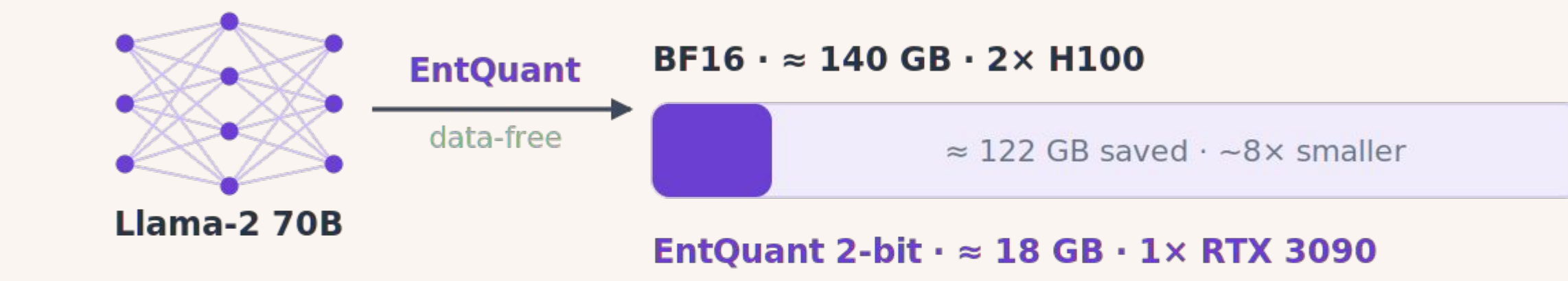
0
calibration data required

Fully data-free.
No training, no fine-tuning,
just model weights.

Any model
including MoEs

Architecture-agnostic.
Works on any new
open-weight model.

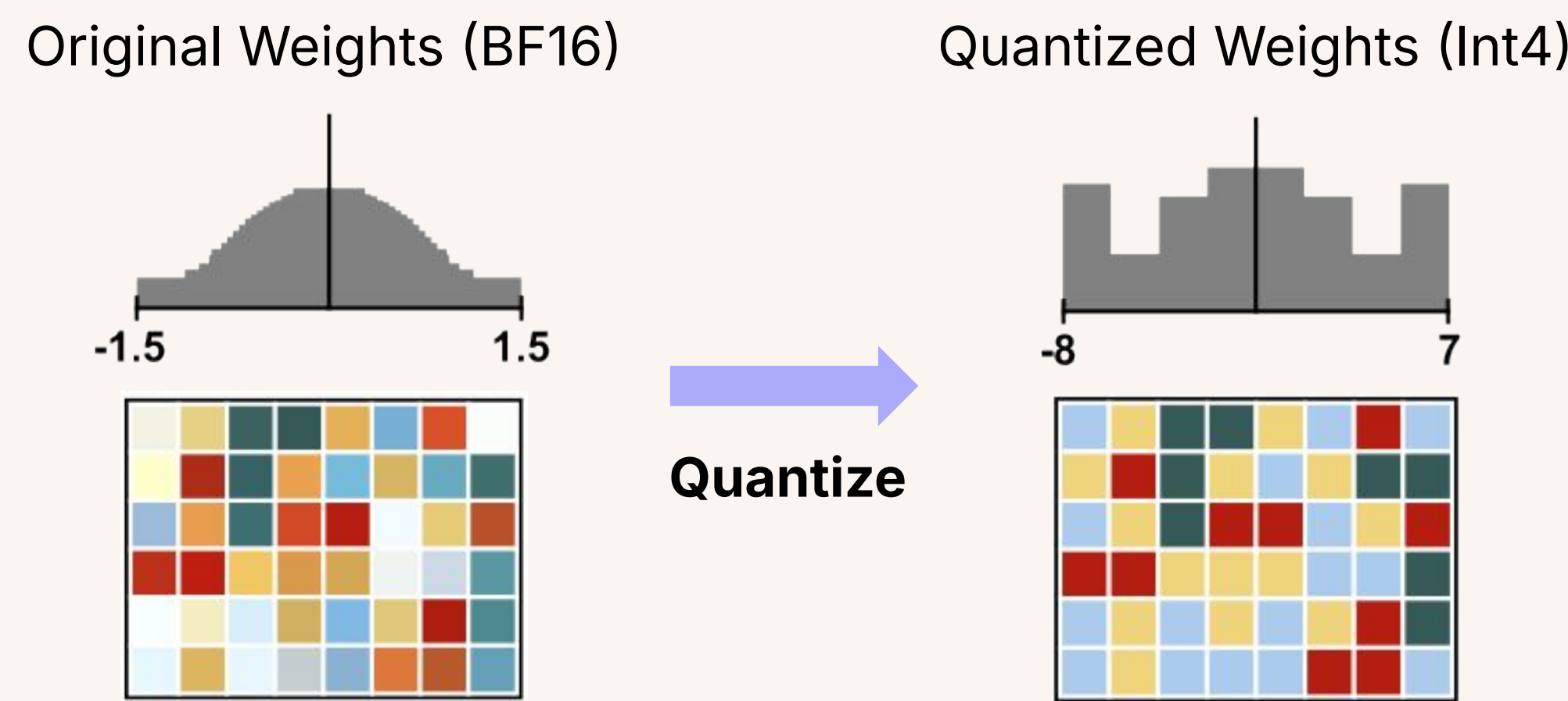
Why Model Compression



WHAT THIS UNLOCKS

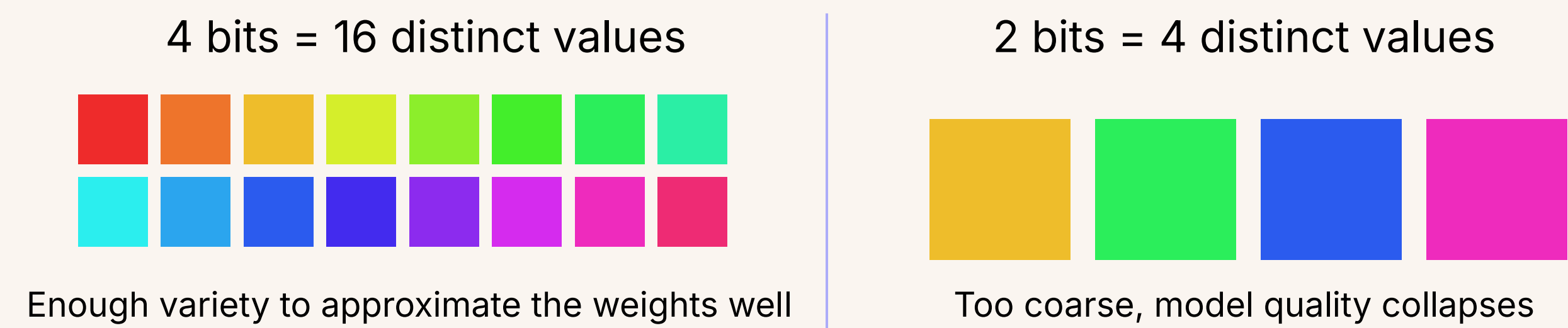
- Self-hosting** on your own hardware
- Overcome resource constraints** one GPU, not a cluster
- Achieve cost savings** consumer-grade GPU

Where Existing Methods Fall Short



Weight quantization works well down to 4 bits ...

... but below, returns start diminishing

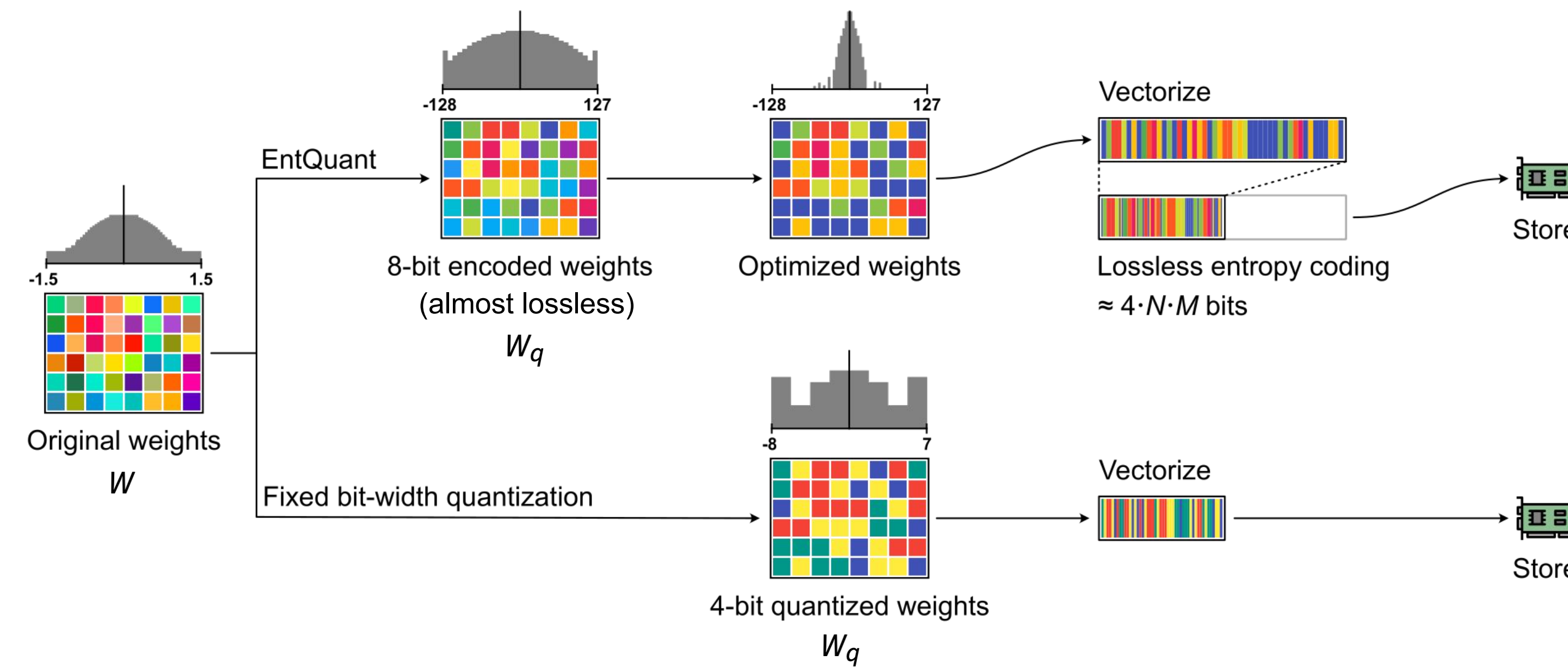


Coupling bit-width to compression rate is problematic in the extreme regime, requiring tricks & hacks to prevent collapse (groups, lattices, recovery training, ...)

👉 Can we go below 4 bits without any data?

EntQuant

Entropy Coding Meets Quantization



Decoupling numerical precision from compression enables more flexibility in the weight optimization

$$\text{Optimized weights} = \arg \min_{W_q} \frac{\|W - \hat{W}\|_1}{\|W\|_1} + \lambda \|W_q\|_1$$

Reconstruction error Surrogate for entropy $H(W_q) = -\sum_{w \in W_q} p(w) \log_2 p(w)$

$W_q = \text{clamp}(\lfloor W/s \rfloor, -Q_{\max}, Q_{\max})$ $s = \max(|W|) / Q_{\max}$

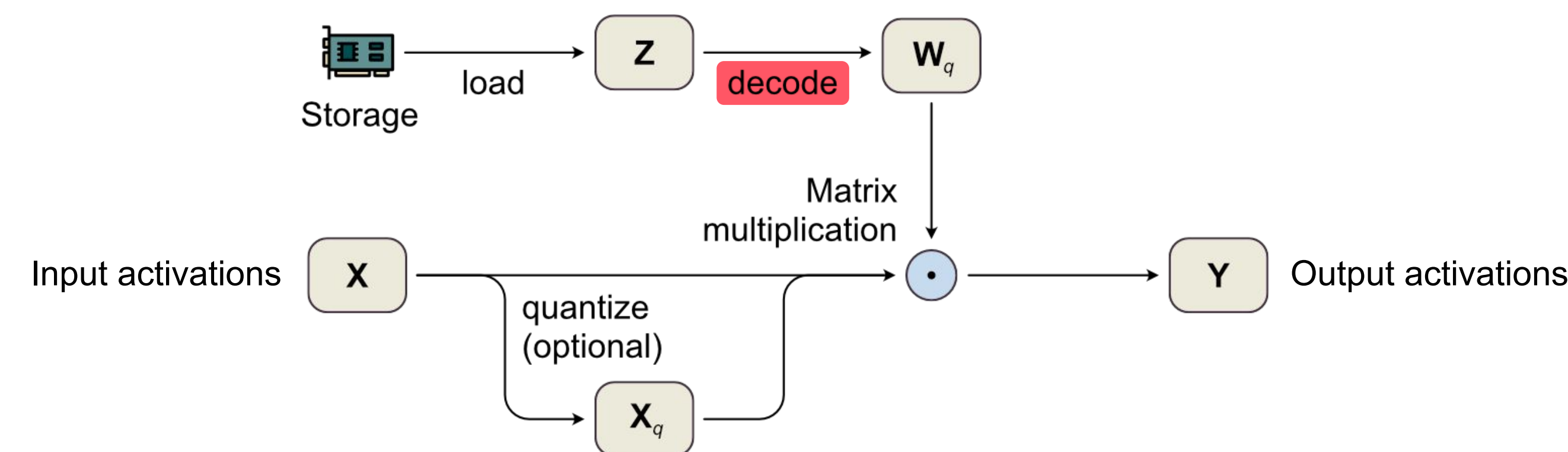
We tweak the scales s per row to minimize entropy

Unique Values (LLaMA-2 7B)

Method	Quantization Levels (bits)		
	4	3	2
Fixed bit-width	16.00	8.00	4.00
EntQuant $\emptyset$	63.89	49.06	34.61

Representing low entropy weights in Float8 / Int8 is much more "expressive" than a comparable fixed bit-width format

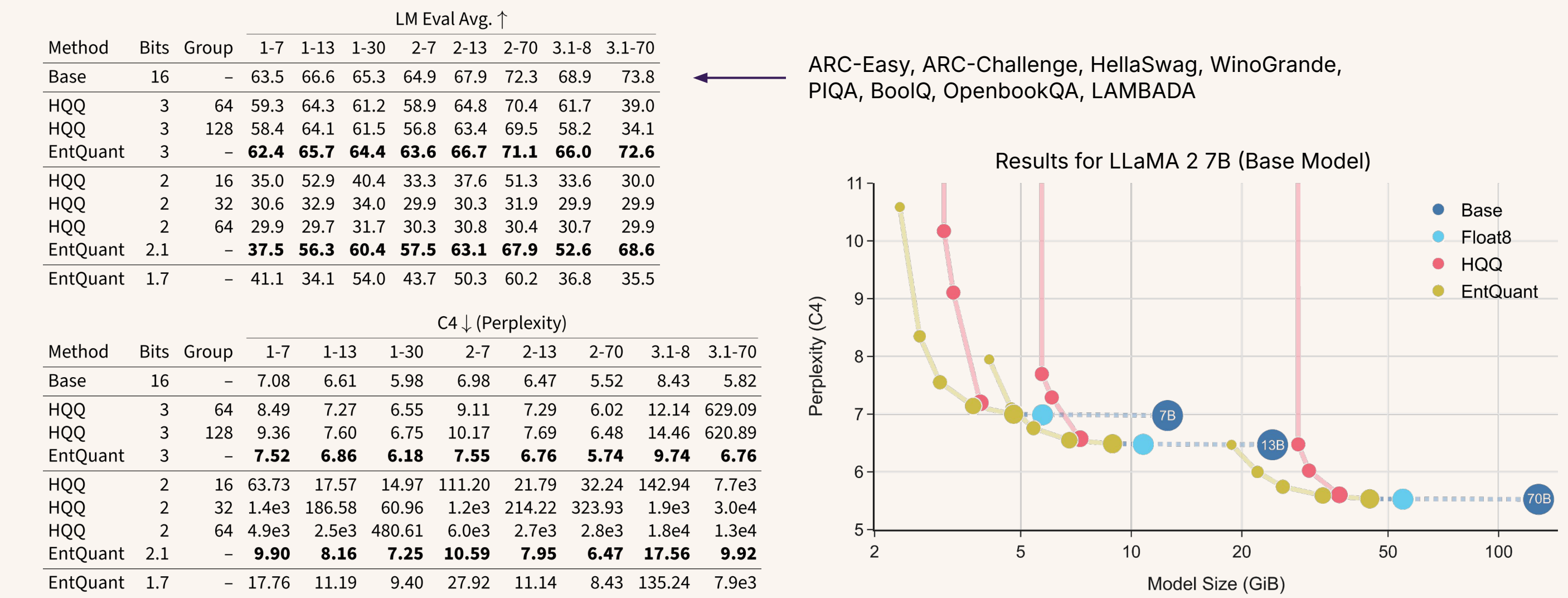
Entropy decoding can be directly built into the inference pipeline



Decoding is the main bottleneck, but highly efficient GPU algorithms exist. With use **ANS** (Asymmetric Numeral Systems), which can achieve a throughput of ~600GB/s in NVIDIA's **nvCOMP** package.

Results

EntQuant outperforms other data-free methods

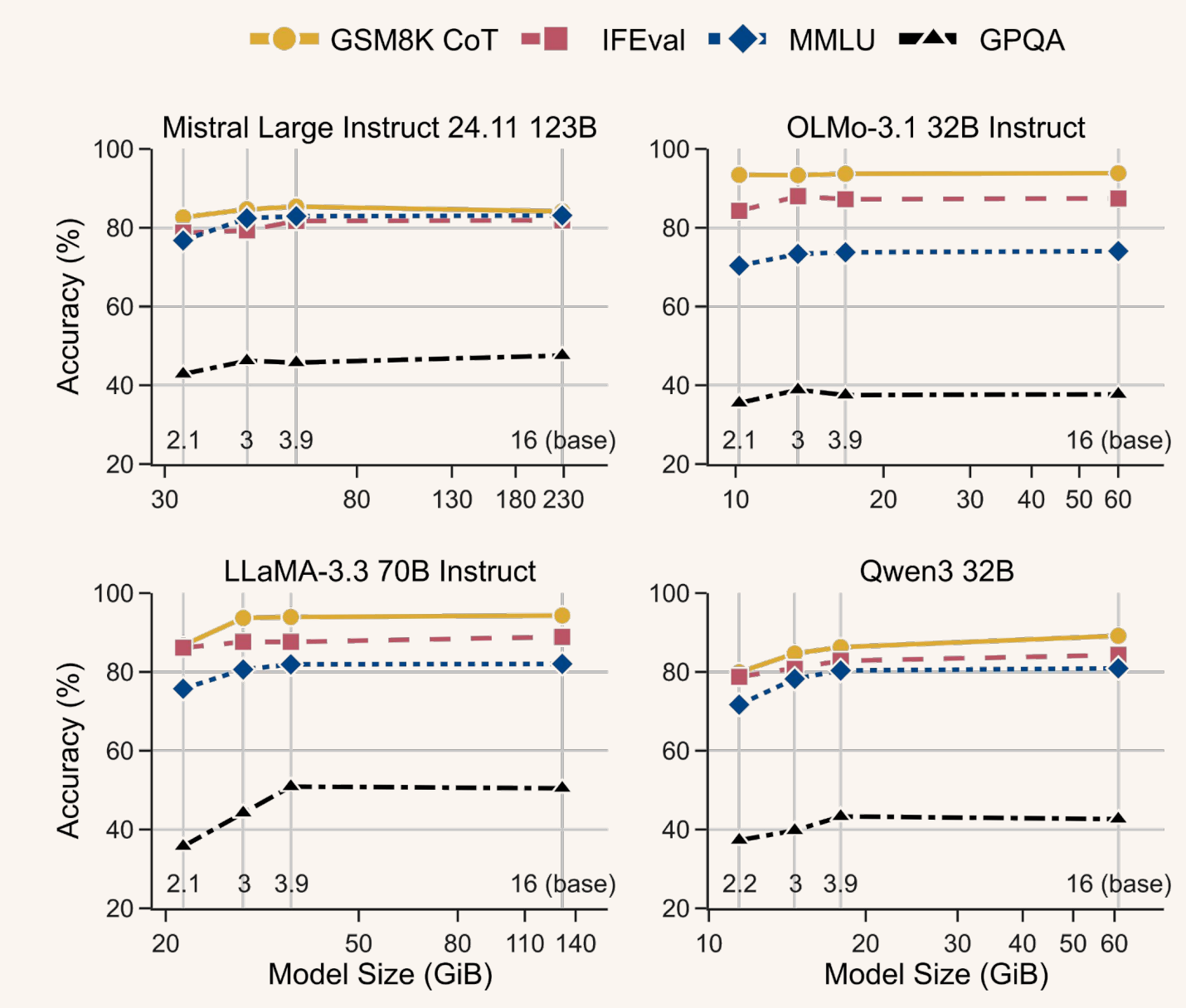


EntQuant keeps up with data-dependent methods ...

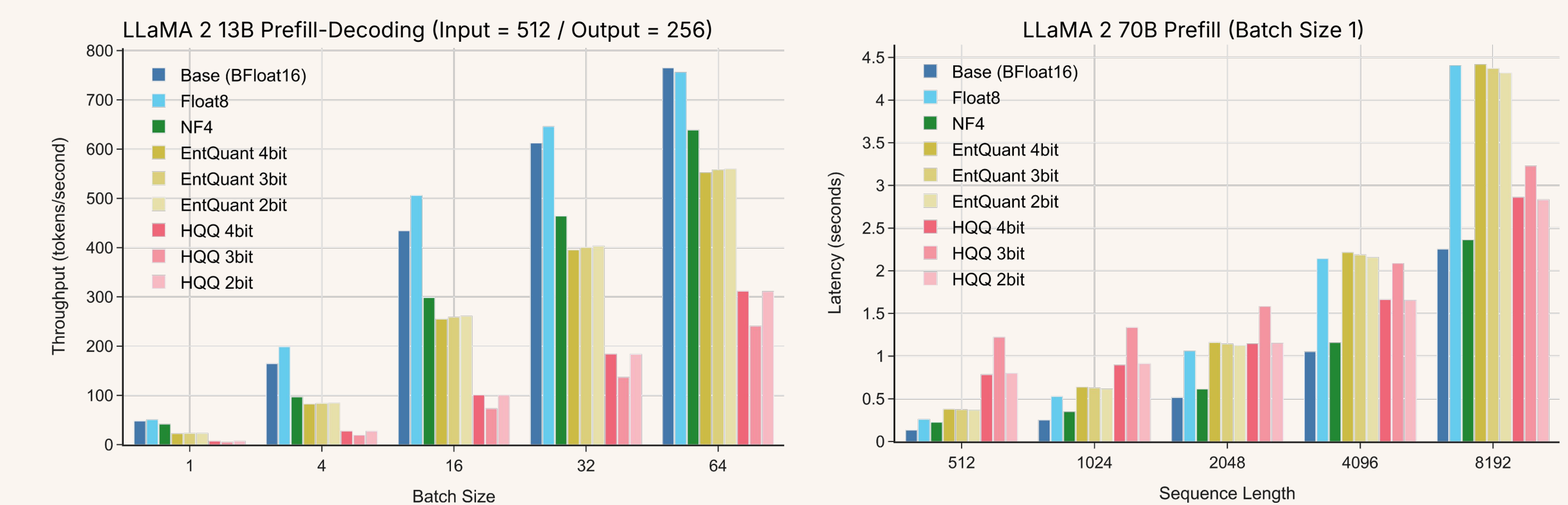
Results for LLaMA 2 70B (Base Model)						
Method	Level	No Calibration	No Training	Compression Runtime	Estimate	
EntQuant	1	✓	✓	<10min (H100)		
GPTQ	2	✗	✓	2-4h (A100-80GiB)		
OmniQuant	2	✗	✓	9-16h (A100-80GiB)		
QuIP#	3	✗	✗	~50h (8× A100-80GiB)		
EfficientQAT	3	✓	✗	~41h (A100-80GiB)		

Method	Bits	Group	C4 $\downarrow$	WikiText-2 $\downarrow$	LM Eval Avg. $\uparrow$
Base	16	-	5.52	3.32	72.8
EntQuant	3	-	5.74	3.62	71.7 (-1.6%)
GPTQ	3	128	5.85	3.85	71.5 (-1.9%)
OmniQuant	3	128	5.85	3.78	71.1 (-2.4%)
QuIP#	3	-	5.67	3.56	72.1 (-0.9%)
EfficientQAT	3	128	5.71	3.61	71.8 (-1.5%)
EntQuant	2.1	-	6.47	4.52	68.6 (-5.8%)
GPTQ	2	128	-	-	34.4 (-52.8%)
OmniQuant	2	128	8.52	6.55	54.9 (-24.6%)
QuIP#	2	-	6.12	4.16	70.9 (-2.6%)
EfficientQAT	2	128	6.48	4.61	68.9 (-5.3%)

... and excels on instruction-tuned models



EntQuant has acceptable inference speed



The entropy decoding overhead is compensated by the Float8 Marlin Kernel