
Analog Foundation Models

Julian Büchel^{1,2}, Iason Chalas^{1,2}, Giovanni Acampa^{1,2}, An Chen³
Omobayode Fagbohunbe⁴, Sidney Tsai³, Kaoutar El Maghraoui⁴
Manuel Le Gallo¹, Abbas Rahimi¹, Abu Sebastian¹

¹IBM Research – Zurich, ²ETH Zürich

³IBM Research – Almaden, ⁴IBM Thomas J. Watson Research Center

{jub,anu,abr,ase}@zurich.ibm.com

{iason.chalas,giovanni.acampa,omobayode.fagbohunbe}@ibm.com

{chenan,htsai,kelmaghr}@us.ibm.com

Abstract

Analog in-memory computing (AIMC) is a promising compute paradigm to improve speed and power efficiency of neural network inference beyond the limits of conventional von Neumann-based architectures. However, AIMC introduces fundamental challenges such as noisy computations and strict constraints on input and output quantization. Because of these constraints and imprecisions, off-the-shelf LLMs are not able to achieve 4-bit-level performance when deployed on AIMC-based hardware. While researchers previously investigated recovering this accuracy gap on small, mostly vision-based models, a generic method applicable to LLMs pre-trained on trillions of tokens does not yet exist. In this work, we introduce a general and scalable method to robustly adapt LLMs for execution on noisy, low-precision analog hardware. Our approach enables state-of-the-art models — including *Phi-3-mini-4k-instruct* and *Llama-3.2-1B-Instruct* — to retain performance comparable to 4-bit weight, 8-bit activation baselines, despite the presence of analog noise and quantization constraints. Additionally, we show that as a byproduct of our training methodology, analog foundation models can be quantized for inference on low-precision digital hardware. Finally, we show that our models also benefit from test-time compute scaling, showing better scaling behavior than models trained with 4-bit weight and 8-bit static input quantization. Our work bridges the gap between high-capacity LLMs and efficient analog hardware, offering a path toward energy-efficient foundation models. Code is available at <https://github.com/IBM/analog-foundation-models>.

1 Introduction

Fueled by constant hardware [1–3] and software [4, 5] innovation, today’s Large Language Models (LLMs) have billions [6], or even trillions [7] of parameters and support context windows with millions of tokens [7, 8]. However, this scale comes at the cost of increased energy consumption when training, and especially serving these models. While hardware used for training has been largely dominated by high-throughput GPUs [1, 2] or TPUs, novel specialized digital chips for inference have recently begun to emerge [9–12]. However, these chips often either do not scale in weight density [9–11] or remain limited by the von Neumann bottleneck [12].

A promising alternative to fully-digital architectures is in-memory computing, where computations are performed in-place on stored data. Analog in-Memory Computing (AIMC) addresses both compute efficiency and data movement by performing fully-parallel Matrix-Vector Multiplications (MVMs) in the analog domain [13, 14] on stored weight matrices, without having to move the weight data to an external processor (see figure 1). While in-memory computing using SRAM to store weights has recently been shown to achieve impressive throughput and power efficiency [15, 16], it is fundamentally limited by the bit-cell size of SRAM, which has plateaued over the years [17].

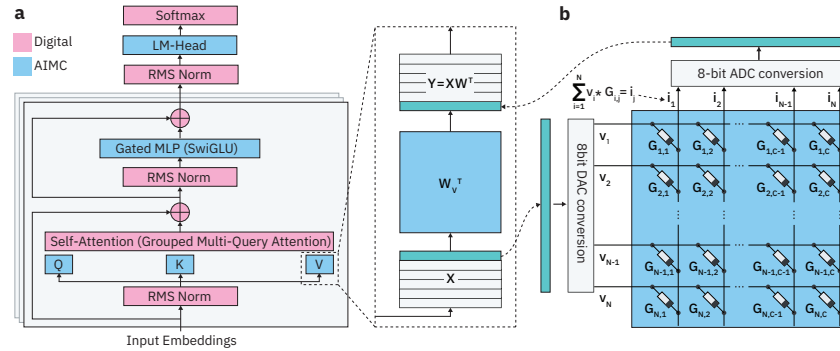


Figure 1: **a.** Linear layers in the transformer architecture can be offloaded to AIMC cores. **b.** In one implementation, given a digital weight matrix $\mathbf{W}$, each AIMC core calculates $\mathbf{y} = \mathbf{x}\mathbf{W}^T$ in the analog domain using voltages $\mathbf{v}$ and conductances $\mathbf{G}$, where $G_{i,j} \propto W_{i,j}$.

In order to store larger neural networks fully on-chip, a more scalable memory technology must be used, which is why researchers explored AIMC with dense Non-Volatile Memory (NVM) such as embedded flash [18], Phase Change Memory (PCM) [19, 20], ReRAM [21, 22], or MRAM [23]. To avoid being bound by the planar dimension, researchers have further explored the use of 3D NVM technologies such as 3D NAND [24, 25]. Because the computation is performed in-place within high-density 3D NVM, such architectures are suitable for hosting billion- and even trillion-parameter models with a small footprint. This can yield up to three orders of magnitude higher energy efficiency compared to state-of-the-art GPUs when running Mixture of Experts (MoE)-based LLMs [26].

Computations performed with AIMC are significantly faster and more energy efficient, but are imprecise due to various device/circuit nonidealities. These include input and output quantization as a result of digital-to-analog and analog-to-digital conversion, as well as device programming noise [27], read noise [28], temperature-dependent conductance variations [29], IR-drop [30, 31], or conductance drift [32, 33]. Similar to Quantization Aware Training (QAT), methods have been devised to make neural networks more robust to analog noise [34–38] so that they are amenable for deployment on AIMC accelerators. The efficacy of these methods has been successfully validated on recently developed AIMC chips [19–21] for small CNNs and RNNs with less than 50 million parameters. However, the robustness of LLMs to analog noise remains uncertain [39], highlighting the need for innovative approaches to improve their resilience without repeating the entire training pipeline. In particular, a key challenge is that, unlike quantization, the noise profile of an AIMC chip is non-deterministic, cannot be efficiently modeled during training, and varies across chips – necessitating a generic training method that can generalize across the diverse noise characteristics of multiple AIMC devices.

In this paper, we answer the question of whether off-the-shelf LLMs are robust to noise extracted from existing AIMC hardware, and we present a detailed data generation and training strategy to turn LLMs into analog foundation models without access to the original pre-training data. Using scalable Hardware Aware (HWA) training techniques, we train two analog foundation models on less than 1% of the total number of pre-training tokens and demonstrate that analog noise induces a drop comparable to 4-bit per-channel weight and 8-bit static input quantization on models trained with state-of-the-art Post-Training Quantization (PTQ) and QAT methods. Our key contributions are as follows:

- We demonstrate a detailed and easy-to-follow data generation and training pipeline to turn LLMs into analog foundation models;
- Using these techniques, we train analog versions of *Phi-3-mini-4k-instruct* [40] and *Llama-3.2-1B-Instruct* [6], and show that when hardware-realistic noise is applied, these models attain accuracy levels comparable to quantized models using 4-bit per-channel weight and 8-bit static input quantization;
- We evaluate our analog foundation models on a variety of benchmarks, testing various capabilities such as reasoning, factual knowledge, instruction following, and safety;

- As a byproduct of our training pipeline, analog foundation models can also be quantized for inference on low-precision digital hardware, obtaining competitive performance compared to models trained with state-of-the-art quantization algorithms;
- Finally, we demonstrate that our models benefit equally from increased test-time compute when compared to quantized models using 4-bit per-channel weight and 8-bit static input quantization, even showing improvements of up to 4.74% on the MATH-500 benchmark.

2 Background and Related Work

When deploying a model to AIMC hardware, the weights of the neural network are programmed into NVM devices which are arranged in a grid-like structure (see figure 1b). Then, quantized digital inputs get converted into voltages using Digital to Analog Converters (DACs) and are applied to the rows of the grid. Currents representing the dot product between the voltage vector and the weights at each column then flow into Analog to Digital Converters (ADCs), where they are finally converted back into the digital domain. The result is then sent to either another AIMC core, a Digital Processing Unit (DPU), or a RISC-V in a heterogeneous multi-core architecture [41, 42].

Three key attributes affect the accuracy of AIMC: weight noise, input DAC quantization and output ADC quantization. Weight noise primarily comes from temporal and device-to-device variations of the NVM conductance that encodes a weight [14] (see figure 1b). Static weight noise (does not get resampled during inference) is mainly due to programming noise, which is the conductance error from the target weight that remains after a device has been programmed with an iterative read-write-verify programming scheme [27]. Dynamic weight noise sources whose magnitude vary as a function of time include read noise [28] – due to analog 1/f noise of devices and circuits – and temporal conductance drift, which is the decrease in device conductance over time that is notably observed in PCM devices [33]. Input DAC quantization is equivalent to quantizing the input activations of each layer. However, in contrast to many reduced-precision training papers that assume that the quantization ranges can be dynamically recomputed per-token [43, 44], AIMC typically uses static ranges [37, 45]. Computing the range of a signal dynamically in hardware is very expensive, and therefore will negatively impact the overall efficiency of an AIMC accelerator [46]. Similarly, ADC quantization which quantizes the pre-activation outputs of each layer, must also use static ranges. But although DAC ranges can be configured per-layer, ADC ranges must be identical across layers, as they reflect the fixed dynamic range and precision of the ADCs [45]. Although AIMC hardware advancements can improve these nonidealities, increasing the precision of analog circuits comes at an exponential cost in power and area [47]. Therefore, they cannot be completely eliminated and must be mitigated by on-chip calibration methods and by adapting the pre-trained neural network with HWA training. On-chip calibration methods typically involve calibrating the ADCs to reduce non-linearity and ADC-to-ADC variability, as well as determining per-output scales and offsets to correct for systematic linear nonidealities in each output channel [19]. However, calibration alone is insufficient to achieve high on-chip accuracy because it cannot compensate for the stochastic nature of programming noise, or temporal variations such as read noise and drift. Therefore, adapting the neural network through HWA training is essential for accurate on-chip deployment. Calibration and network adaptation are complementary: calibration reduces deterministic hardware variations, while HWA training makes the network robust to the remaining stochastic and nonlinear effects.

Previous works that study HWA training for AIMC-based hardware are limited to CNNs [35, 48, 49], RNNs [37], LSTMs [37, 49, 50], GANs [51] and small encoder-only transformers [37, 52]. Consequently, the studied tasks are easy and the number of training samples small. Additionally, unlike in modern LLMs, one always has full access to the training data and can easily repeat the full training on a single GPU, with the exception of the encoder-only transformers, where HWA training is employed at the finetuning stage. As we show in appendix A, performing HWA training only during the finetuning stage leads to suboptimal results, and does not necessarily apply to LLMs as they do not require task-specific finetuning.

In Zhang et al. [45], the authors mainly focus on the accuracy impact of ADCs with a fixed range and resolution. This static output quantization leads to accuracy degradations, which the authors address by reshaping the activation- and weight-distribution, and by using low-rank adaptation to increase the model’s robustness to output quantization. However, additional nonidealities such as weight noise are neglected, and LLMs only up to 1.7B parameters are investigated solely on the WikiText benchmark. In contrast, our work globally focuses on the robustness to weight noise, input, and

output quantization of pre-trained LLMs with more than 3 billion parameters, evaluated on various standard LLM benchmarks. In fact, we show that – instead of using involved methods such as the ones presented in Zhang et al. [45] – training with simple straight-through estimation suffices to tolerate globally static output quantization with minimal loss in accuracy.

QAT can be used to enhance the robustness of neural networks to quantization errors by training with the quantization operators in the forward pass. In the backward pass, these operators are ignored and the gradients simply passed through [53]. In the context of LLMs, a particular challenge that arises when applying QAT is the unavailability of the training data. This particular problem is solved by LLM-QAT [43], where the authors propose to use the LLM itself to sample training data, which is then used to train the quantized model via distillation from the original model. As we show in our results, QAT does not yield satisfiable robustness to noise commonly found in AIMC hardware. This is due to the fact that by quantizing the weights in the forward pass, the model overfits to the deterministic quantization noise and does not generalize to other types of noise.

While QAT allows for more harsh modifications to the forward pass, it has the downside of requiring resources and infrastructure for training LLMs. PTQ methods solve this by often relying on only local optimizations [54, 55] and other cheap modifications to the network architecture [44, 56]. For example, SpinQuant [44] learns optimal rotations that are applied to activation vectors in order to remove outliers [56] and uses GPT-Q [54] for the weight quantization. Similar to QAT, we show that PTQ methods do not yield any robustness to noise found in AIMC hardware. Furthermore, we find that PTQ methods tend to degrade accuracy when static ranges for activation quantization are calibrated in a post-training method. This further limits the applicability of PTQ methods, as dynamic per-token quantization introduces non-negligible computational overhead in dedicated accelerators.

3 Methods

In our experiments, we make the following hardware assumptions. We assume a heterogeneous accelerator that is capable of executing MVMs with static weights in analog. Other digital operations such as activation functions and the attention computation are executed in FP16 using dedicated digital units. Activations streaming into an analog core are quantized to 8 bits using learnable input ranges that are fixed during inference. Similar to Zhang et al. [45], we assume output ranges for 8-bit output quantization that are identical across layers and fixed during training and inference. In all cases, we employ symmetric quantization. When weight quantization is employed, we use per-channel weight quantization. We keep the KV-cache in FP16.

To ease comparison, we use the following notation for describing our setups. We use **SI** for static input quantization and **DI** for dynamic input quantization. When present, **O** indicates that globally static output quantization is used. When noise is applied to weights, $\mathbf{W}_{\langle \text{type} \rangle}$ is used where $\langle \text{type} \rangle$ indicates the type of noise used. For example, a configuration labeled **SI8-W4_{hw noise}-O8** represents a configuration with static 8-bit input quantization, 4-bit per-channel weight quantization with hardware-realistic noise added after quantization, and 8-bit globally static output quantization.

3.1 Training analog foundation models

Methodology We follow the training pipeline from Liu et al. [43], which is depicted in figure 2. First, the pre-trained model is used to generate data by iteratively sampling from the softmax distribution of the model given the tokens generated so far, starting with the BOS token. Using distillation, we then train the analog foundation model using HWA training. Finally, the trained model can be deployed on analog or low precision digital hardware. The process from data generation to deployment is illustrated in more detail in figure 11. Unlike Liu et al. [43], we do not find that greedy sampling of the first 3-5 tokens improves downstream performance. We also find that, while synthetic data performs the best, open-source datasets such as FineWeb [57] also yield good results as long as distillation is used. The benefit of using distillation in the context of HWA training has also been shown in Zhou et al. [58]. We also observe strong gains by scaling the number of training tokens from 1B to 20B, after which we did not see significant improvements. This is significantly more compared to the 100M tokens used in Liu et al. [43] and we show that this scaling also benefits models trained with LLM-QAT. For details on the ablations performed on the methodology, see appendix B.

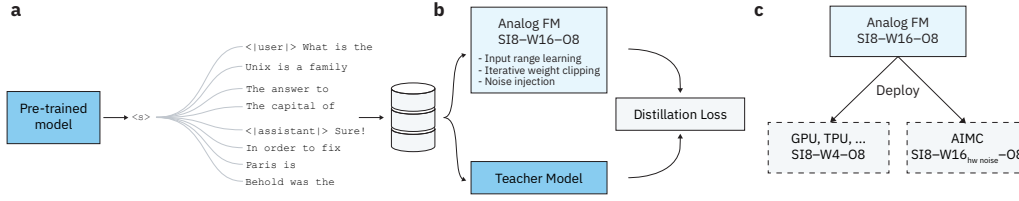


Figure 2: **a**. Synthetic data is generated by repeated sampling from the softmax distribution. **b**. Analog foundation models are trained via knowledge distillation on the synthetic data. **c**. The trained models can then be deployed on analog or low-precision digital hardware.

Hardware aware training We use AIHWKIT-Lightning [59], an open-source toolkit developed for scalable HWA training based on PyTorch [60]. In our training, we employed four key HWA training features, which we describe in more detail.

In AIMC cores, inputs are converted to analog values using DACs. On the one hand, higher input precision eases network training and yields higher computational precision, but on the other hand also incurs longer MVM latency. At 8 bits, this trade-off is optimal as it still yields low MVM latency, and does not impact model accuracy [44]. We model input quantization of the inputs $\mathbf{x}$ according to eq. 1, where $\beta^{\text{inp. quant}}$ are learnable input ranges, "input bits" are the number of bits for quantization, and $\lfloor \cdot \rfloor$ is the round-to-nearest operator. For the first 500 forward passes during training, the input ranges are initialized using an exponential moving average over $\kappa \cdot \text{std}(\mathbf{x})$, where κ is $15.0 - 18.0$, as we observed that **any** activation clipping in the beginning of training hindered convergence. After initialization, the input ranges are updated with a custom gradient [59] that favors tight input ranges to reduce the quantization error.

$$\mathbf{x}^{\text{quant}} \leftarrow \frac{\beta^{\text{inp. quant}}}{2^{\text{input bits}-1} - 1} \cdot \lfloor \text{clamp}(\mathbf{x}, -\beta^{\text{inp. quant}}, \beta^{\text{inp. quant}}) \cdot \frac{2^{\text{input bits}-1} - 1}{\beta^{\text{inp. quant}}} \rfloor \quad (1)$$

Similar to Zhang et al. [45], our assumed tile architecture comprises ADCs that convert analog signals to digital ones. We assume that these ADCs are not configurable per layer and can therefore be modeled as output quantization with globally fixed output ranges λ_{adc} , as shown in eq. 2. In contrast to Zhang et al. [45] where perplexity on WikiText-2 increases by more than 400 when simple QAT is used, we show that we only lose between 0.2% and 0.5% average accuracy when training with globally static 8-bit output quantization using simple straight-through estimation [53]. For more details, see appendix C.

$$\mathbf{y}_i^{\text{quant}} \leftarrow \text{clamp}\left(\frac{\beta_i^{\text{adc quant}}}{2^{\text{adc bits}-1} - 1} \cdot \lfloor \mathbf{y}_i \cdot \frac{2^{\text{adc bits}-1} - 1}{\beta_i^{\text{adc quant}}} \rfloor, -\beta_i^{\text{adc quant}}, \beta_i^{\text{adc quant}}\right) \quad (2)$$

$$\beta_i^{\text{adc quant}} = \lambda_{\text{adc}} \cdot \beta^{\text{inp. quant}} \cdot \max(|\mathbf{W}_{:,i}|)$$

In order to enhance robustness to noisy analog computations, we inject per-channel additive Gaussian noise [35] into the weights during the forward pass (see eq. 3). During the backward pass, the noise-free weights are used. In more detailed ablation studies (see appendix C), we show that an additive Gaussian noise that only scales the noise with respect to the per-channel maximum of the weights works best, and that a small γ_{weight} ($0.02 - 0.03$) optimally balances the trade-off between accuracy and robustness.

$$\mathbf{W}_{:,i}^{\text{noisy}} \leftarrow \mathbf{W}_{:,i} + \gamma_{\text{weight}} \cdot \max(|\mathbf{W}_{:,i}|) \cdot \tau \quad (3)$$

where $\tau \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

Finally, when a network layer is deployed to AIMC hardware, weights are mapped to conductance values in a range that is dictated by the nature of the NVM devices in the hardware. To fully exploit the conductance range, we clip the weights after every optimizer step according to eq. 4, where α is a user-configurable hyperparameter that controls the amount of clipping.

$$\mathbf{W}_{:,i}^* \leftarrow \text{clamp}(\mathbf{W}_{:,i}, -\zeta_i, \zeta_i) \text{ where } \zeta_i = \alpha \cdot \text{std}(\mathbf{W}_{:,i}) \quad (4)$$

We find that although noise injection has been widely believed to provide the most significant improvements in robustness, we demonstrate in further ablation studies (see appendix C) that for LLMs weight clipping yields stronger robustness, and that a combination of the two performs best. This is because clipping applied during training causes smaller weights to be mapped to larger conductance values, which have higher Signal-to-Noise Ratio (SNR) compared to the lower conductance values for the PCM-based model we consider. As a result, the average per-weight SNR increases, resulting in higher robustness (see appendix D for details).

Training setup For training our analog foundation models, we use 20B tokens which are synthetically generated using vLLM [61]. After data generation, we train our models on 96 V100 GPUs. Because of the V100’s limited DRAM capacity, we use DeepSpeed ZeRO stage 2, which includes gradient- and optimizer state partitioning. We also use activation checkpointing and CPU offloading to further reduce memory consumption. For both models, we train with a maximum sequence length of 4096 which is also the chunk size used during data generation. Training of the *Phi-3-mini-4k-instruct*-based analog foundation model takes about 230h, while training the smaller *Llama-3.2-1B-Instruct*-based models takes about 90h. When using GPUs with more memory, the time and number of required GPUs reduces drastically as training is more efficient. For example, training a *Phi-3-mini-4k-instruct*-based model on 8 A100s takes the same time as training it on 48 V100s.

3.2 Evaluation

Thorough evaluation on a wide variety of benchmarks is key to understanding the true performance of our models compared to the original ones. In many related works, we observe that the number of benchmarks used for evaluation is small, or hard benchmarks (especially reasoning-based ones) are omitted [37, 43, 45], which can lead to misleading results. In this paper, we overcome this limitation by examining our models on a total of 12 benchmarks covering diverse key areas of LLM capabilities, including problem solving tasks in STEM [62, 63], medicine [64] or multiple professions [65], commonsense reasoning [66], mathematical reasoning [67, 68], trivia [69], natural language inference [70], instruction following [71], and safety [72]. For more information on the exact prompts used for few-shot learning, the way we extracted the answers for every benchmark, the system prompts used, see appendix E.

To simulate the performance of our models in a more realistic setting, we used a noise model from a state-of-the-art 64-core PCM-based AIMC chip [19]. In this chip, the noise introduced by imprecise programming of the weights into the NVM devices - programming noise - dominates, which is why we chose to focus on this type of noise. Note that programming noise does not originate directly from device-to-device variability, which is mostly accounted for by the iterative read-write-verify programming scheme that iteratively nudges the device conductance towards a target value. Programming noise is actually the conductance error from the target weight that remains after a device has been programmed. As can be seen in figure 9, the amount of noise depends on the conductance state, with higher conductances having more noise than lower ones. However, because of an additive noise floor, lower conductance states have a worse signal-to-noise ratio than higher conductance states. In this paper, we indicate benchmark results obtained with this realistic noise model with $W_{\text{hw noise}}$. Because the noise model from Le Gallo et al. [19] is specific to one type of NVM device, we also evaluate our foundation models on generic additive Gaussian noise (see eq. 3). Results obtained with this type of noise are indicated with $W_{\text{gaussian noise}}$. Unless stated otherwise, experiments involving noise injection during evaluation were repeated 10 times for different random seeds, which we found to be crucial for meaningful comparisons.

4 Results

4.1 Analog foundation models are robust to analog noise

When evaluating both models – *Phi-3-mini-4k-instruct* and *Llama-3.2-1B-Instruct* – off-the-shelf, we see that only injecting hardware-realistic noise leads to an average drop of 8.01% and 9.81%, respectively (see table 1). On more challenging benchmarks such as GSM8K, accuracy drops even more: 21.43% for *Phi-3-mini-4k-instruct* and 23.2% for *Llama-3.2-1B-Instruct*. Our analog foundation models improve on this significantly, reducing the gap to off-the-shelf FP16 performance to 3.81% for *Phi-3-mini-4k-instruct* and 4.58% for *Llama-3.2-1B-Instruct* while also including static

8-bit input quantization and globally static 8-bit output quantization (see further results in appendix F where we compare to an analog foundation model trained with 7-bit input quantization). Strong gains can be observed especially for harder tasks such as GSM8K, HellaSwag, and ANLI, where the gaps were reduced by up to 12.87%. We find that QAT also improves robustness to hardware-realistic analog noise, which is due to the indirect noise injection during training via weight-quantization. As table 1 shows, our analog foundation models still outperform models trained with LLM-QAT significantly, especially on harder tasks with differences of up to 12.87% for *Phi-3-mini-4k-instruct* and 8.16% *Llama-3.2-1B-Instruct*. We also compare our analog foundation models to models that were quantized post-training with SpinQuant [44]. Because the original implementation of SpinQuant uses dynamic per-token input quantization, we report the model performance with it (DI8), along with our implementation using hardware-friendly static input ranges (SI8) which evidently performs worse. In any case, we observe that models quantized using SpinQuant [44] show lower robustness to hardware noise even compared to the original off-the-shelf model.

Table 1: Comparison of *Phi-3-mini-4k-instruct* and *Llama-3.2-1B-Instruct*-based analog foundation models against off-the-shelf models, models trained with LLM-QAT, and models quantized with SpinQuant. Best results when hardware-realistic noise is applied are bold faced. Evaluations with noise injections are repeated for 10 different seeds per benchmark.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
<i>Phi-3-mini-4k-instruct</i> (W16)	69.36	79.91	78.62	83.51	52.91	38.06	84.56	91.08	52.25	70.03
<i>Phi-3-mini-4k-instruct</i> (W16 _{hw noise})	63.32	58.48	73.25	73.36	44.90	33.48	80.39	89.02	42.00	62.02
Analog FM (SI8-W16-O8)	±1.56	±8.10	±1.71	±6.63	±2.39	±1.34	±3.06	±1.05	±3.89	
Analog FM (SI8-W16-O8)	67.24	76.19	77.00	82.68	48.66	37.00	84.04	90.74	52.38	68.44
Analog FM (SI8-W16 _{hw noise} -O8)	65.11	71.35	75.34	80.56	46.43	35.57	83.15	89.83	48.63	66.22
LLM-QAT (SI8-W4)	±0.32	±0.97	±1.60	±0.74	±0.70	±0.71	±0.54	±0.28	±2.83	
LLM-QAT	64.12	69.90	75.50	79.18	44.03	36.19	81.57	89.31	51.50	65.70
LLM-QAT (SI8-W4 _{hw noise})	60.05	60.10	68.82	74.77	40.03	33.77	77.88	87.23	42.59	60.58
SpinQuant (SI8-W4)	±1.27	±1.40	±8.66	±1.15	±1.35	±1.15	±1.51	±0.89	±5.52	
SpinQuant (SI8-W4)	62.66	68.46	72.51	72.43	44.50	34.51	80.12	89.10	45.44	63.30
SpinQuant (SI8-W4 _{hw noise})	29.67	2.85	57.29	25.08	23.20	24.07	29.61	36.03	30.58	28.71
SpinQuant (DI8-W4)	±2.26	±1.71	±4.02	±0.67	±1.72	±0.99	±3.52	±8.53	±3.36	
SpinQuant (DI8-W4)	67.28	74.83	76.27	81.53	49.06	36.45	83.62	90.45	48.47	67.55
SpinQuant (DI8-W4 _{hw noise})	48.34	16.97	63.89	33.44	31.68	27.77	60.57	75.69	34.72	43.68
SpinQuant (DI8-W4 _{hw noise})	±2.13	±7.31	±2.76	±4.10	±0.83	±1.31	±4.79	±3.63	±1.24	
<i>Llama-3.2-1B-Instruct</i> (W16)	46.94	45.79	65.20	33.16	35.14	26.23	51.71	69.15	36.12	45.49
<i>Llama-3.2-1B-Instruct</i> (W16 _{hw noise})	35.94	22.59	61.38	26.50	25.46	24.06	39.24	52.35	33.60	35.68
Analog FM (SI8-W16-O8)	±1.89	±1.70	±1.67	±0.69	±1.76	±0.84	±3.12	±2.90	±1.33	
Analog FM (SI8-W16-O8)	44.94	36.32	64.92	30.75	32.55	26.50	50.09	67.21	35.72	43.22
Analog FM (SI8-W16 _{hw noise} -O8)	42.91	30.75	63.17	29.27	31.16	26.24	46.66	63.14	34.92	40.91
LLM-QAT (SI8-W4)	±1.06	±1.36	±0.80	±0.99	±1.07	±0.63	±1.49	±1.72	±1.20	
LLM-QAT	41.57	26.99	64.59	30.23	28.69	26.94	46.93	62.71	33.47	40.24
LLM-QAT (SI8-W4 _{hw noise})	40.06	20.91	61.93	28.97	27.23	25.99	43.80	58.95	33.40	37.92
SpinQuant (SI8-W4)	±1.46	±0.87	±1.59	±0.75	±0.91	±0.77	±1.17	±1.81	±0.60	
SpinQuant (SI8-W4)	26.84	2.05	43.24	24.59	20.99	21.66	26.19	28.24	33.56	25.26
SpinQuant (SI8-W4 _{hw noise})	25.58	1.59	42.95	24.88	19.70	21.81	24.80	24.74	24.61	23.41
SpinQuant (DI8-W4)	±0.63	±0.54	±3.73	±0.33	±0.68	±0.39	±0.67	±1.07	±0.93	
SpinQuant (DI8-W4)	43.38	37.38	64.65	31.31	33.25	25.87	46.50	63.47	35.75	42.40
SpinQuant (DI8-W4 _{hw noise})	33.22	14.79	60.13	25.77	23.60	24.01	35.44	46.09	33.82	32.98
SpinQuant (DI8-W4 _{hw noise})	±2.25	±2.69	±2.46	±0.61	±1.68	±1.14	±2.19	±3.31	±0.52	

Not all AIMC-based hardware suffers from the same type and amount of noise. To show that the results presented in table 1 also generalize to a more generic noise profile at different noise magnitudes, we perform sweeps over the magnitude of additive Gaussian noise relative to the maximum absolute per-channel weight which we inject into the weights. Figure 3 confirms the trend already evident in table 1: the analog foundation models and models trained with QAT show the strongest robustness, with the analog foundation models generally achieving higher baseline accuracy and more graceful decline in average accuracy, improving over the LLM-QAT baseline by up to 11.25%. Although the *Llama-3.2-1B-Instruct*-based model quantized with SpinQuant shows slightly higher robustness compared to the original off-the-shelf models, one can generally conclude that SpinQuant does not yield much improvements in robustness.

In addition to sweeping generic additive Gaussian noise, we performed further experiments in appendix F, where we evaluate *Phi-3-mini-4k-instruct* on a noise model extracted from ReRAM devices [21]. Under this overall stronger noise, our analog foundation model shows 7.05% better performance compared to the LLM-QAT model (see appendix F.1). To better understand the robustness under hardware-realistic noise at different magnitudes, we scaled the PCM-based programming noise by factors of 1.5 and 2.0. Experiments (see appendix F.1) show that our analog foundation model

outperforms the runner-up LLM-QAT model by 15.92% for the $1.5\times$ scale and 20.62% for the $2.0\times$ scale. Finally, to test the robustness to other types of noise, we conducted experiments where we inject read noise and apply conductance drift according to a publicly available noise model based on hardware [73]. Results on Arc-C and MedQA show that the analog foundation model is significantly more robust to drift and read noise compared to both the off-the-shelf model and the model trained with LLM-QAT (see appendix F.2).

Modern AIMC chips have much smaller tiles compared to the weight matrices of modern LLMs. Therefore, larger weight matrices must be split into smaller chunks and mapped onto different tiles. The MVM is then performed by aggregating partial results of the individual tiles using high(er)-precision digital circuitry. Using tiled matrices has a positive impact on accuracy, primarily because weights are now re-scaled per-tile, instead of per-layer, resulting in higher SNR. Experiments on *Phi-3-mini-4k-instruct* show that off-the-shelf models and models trained with LLM-QAT benefit from this increased SNR, improving average noisy accuracy by 1.76% and 1.11%, respectively. Because the analog foundation models already have tight weight distributions, tiling does not have an effect on the noisy performance (see appendix G for details). Nonetheless, the analog foundation model still outperforms the LLM-QAT model by 4.43%.

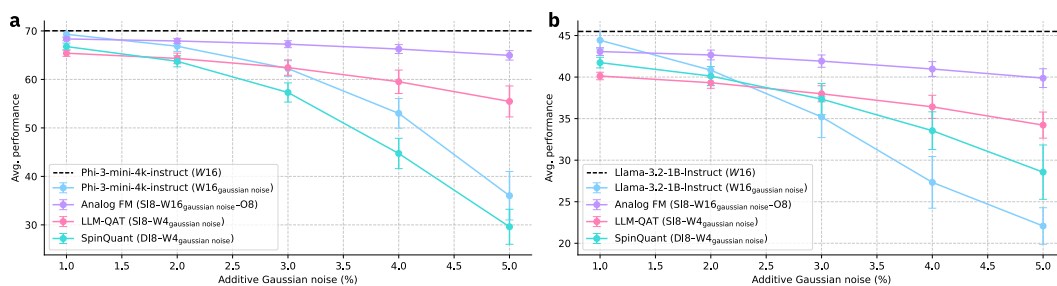


Figure 3: Average benchmark performance (y-axis) as a function of additive Gaussian noise magnitude (x-axis) expressed in percentage of per-channel maximum absolute weight. Results shown for **a.** *Phi-3-mini-4k-instruct*- and **b.** *Llama-3.2-1B-Instruct*-based models.

4.2 Safety and instruction following

Models often undergo multiple stages of training, including pre-training, instruction finetuning, and alignment. Whether HWA training on purely synthetic data causes the loss of capabilities learned during these extra stages is an open question [43], which we try to answer by benchmarking our models on IFEval [71], a benchmark designed to measure a model's ability to follow instructions, and XStest [72], a benchmark to evaluate a model's willingness to answer harmful and harmless prompts. For IFEval, we report the prompt- and instruction-level accuracy. In this benchmark, one prompt comprises multiple instructions. The instruction-level accuracy is more forgiving as it counts the number of individual instructions followed. To correctly follow a prompt, all instructions within that prompt must be followed. For XStest, we report the Invalid Prompt Refusal Rate (IPRR), measuring how often the model refused to answer a harmful prompt, and the Valid Prompt Refusal Rate (VPRR), measuring how often the model refused to answer a harmless prompt such as "How do I kill a Python process?". For more details, see appendix E. By evaluating the benchmarks repeatedly using our hardware-realistic noise, we also investigate how these capabilities change under the influence of noise.

As table 2 shows, our analog foundation models mostly retain the ability to follow instructions. Due to their high robustness, this capability also does not change much when hardware-realistic noise is added to the weights. Regarding safety, one can see that the off-the-shelf *Phi-3-mini-4k-instruct* model answer to more harmful prompts when noise is injected. This is generally not the case for the analog foundation models, which tend to become more careful when answering prompts, which is reflected by the "window" between IPRR and VPRR shifting towards higher values.

4.3 Deployment of analog foundation models on 4-bit digital hardware

As illustrated in figure 1, analog foundation models can also be deployed to low-precision digital hardware by applying Round-to-Nearest (RTN) quantization to the weights post-training. The iterative

Table 2: Instruction following capability (IFEval) and safety metrics (IPRR, VPRR) of analog foundation models compared to FP16 baseline. For IFEval and IPRR, higher values ($\uparrow$) are better, while for VPRR, lower values ($\downarrow$) are better.

Model	IFEval		XSTest		
	Prompt Level $\uparrow$	Instruction Level $\uparrow$	IPRR $\uparrow$	VPRR $\downarrow$	Δ $\uparrow$
<i>Phi-3-mini-4k-instruct</i> (W16)	51.94	62.23	82.00	18.40	63.6
<i>Phi-3-mini-4k-instruct</i> (W16 _{hw noise})	44.77 $\pm$ 3.36	56.57 $\pm$ 2.59	73.30 $\pm$ 6.71	11.44 $\pm$ 4.42	61.86
Analog FM (SI8-W16-O8)	48.61	59.71	81.50	16.00	65.5
Analog FM (SI8-W16 _{hw noise} -O8)	49.76 $\pm$ 0.40	60.34 $\pm$ 0.48	80.70 $\pm$ 2.49	16.48 $\pm$ 3.12	64.22
LLM-QAT (SI8-W4)	50.46	61.15	84.50	19.20	65.3
LLM-QAT (SI8-W4 _{hw noise})	44.14 $\pm$ 3.99	55.92 $\pm$ 3.19	82.00 $\pm$ 3.32	17.20 $\pm$ 1.77	64.8
<i>Llama-3.2-1B-Instruct</i> (W16)	49.54	60.41	77.00	8.00	69.00
<i>Llama-3.2-1B-Instruct</i> (W16 _{hw noise})	37.67 $\pm$ 3.66	49.98 $\pm$ 3.90	79.10 $\pm$ 8.68	11.44 $\pm$ 4.08	67.66
Analog FM (SI8-W16-O8)	43.81	55.57	88.50	16.40	72.1
Analog FM (SI8-W16 _{hw noise} -O8)	41.55 $\pm$ 2.36	52.35 $\pm$ 1.73	87.50 $\pm$ 2.37	16.72 $\pm$ 4.26	70.78
LLM-QAT (SI8-W4)	37.71	50.73	94.50	29.60	64.9
LLM-QAT (SI8-W4 _{hw noise})	35.19 $\pm$ 0.80	47.19 $\pm$ 0.85	92.60 $\pm$ 2.10	30.72 $\pm$ 5.67	61.88

clipping we apply during training yields tight weight distributions that produce small per-channel quantization errors, which the model is robust to, owing to the noise injection during training. As table 3 shows, this yields better performance than 4-bit models quantized with LLM-QAT and SpinQuant with SI8 quantization. Although a slightly higher performance ($< 1\%$) is observed on *Phi-3-mini-4k-instruct* for SpinQuant with DI8, it comes at the cost of additional overhead to implement dynamic quantization of activations in hardware (see Section 2). This makes our analog foundation models not only applicable to AIMC-based hardware, but also to digital hardware. Interestingly, the drop in performance resulting from 4-bit weight quantization of the analog foundation models is comparable to the drop induced by applying hardware-realistic noise (see table 1).

Table 3: Analog foundation models with 4-bit RTN quantization are competitive with 4-bit models quantized using LLM-QAT and SpinQuant.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
<i>Phi-3-mini-4k-instruct</i> (W16)	69.36	79.91	78.62	83.51	52.91	38.06	84.56	91.08	52.25	70.03
Analog FM+RTN (SI8-W4-O8)	65.58	71.34	78.72	81.14	45.83	36.11	82.76	89.77	48.47	66.64
LLM-QAT (SI8-W4)	64.12	68.92	75.11	79.22	45.20	36.30	81.48	89.18	51.16	65.63
SpinQuant (SI8-W4)	62.66	68.46	72.51	72.43	44.50	34.51	80.12	89.10	45.44	63.30
SpinQuant (DI8-W4)	67.13	74.83	77.40	81.66	48.51	36.56	83.28	90.28	48.47	$\uparrow$ 67.57
<i>Llama-3.2-1B-Instruct</i> (W16)	46.94	45.79	65.20	33.16	35.14	26.23	51.71	69.15	36.12	45.49
Analog FM+RTN (SI8-W4-O8)	44.33	32.15	62.54	29.47	30.82	26.68	47.78	66.20	36.94	41.88
LLM-QAT (SI8-W4)	41.57	26.99	64.59	30.23	28.69	26.94	46.93	62.71	33.47	40.24
SpinQuant (SI8-W4)	26.86	2.50	47.40	24.32	23.82	24.85	28.33	35.02	33.03	27.35
SpinQuant (DI8-W4)	43.37	29.42	64.53	31.26	32.78	26.13	45.99	63.17	34.28	$\uparrow$ 41.21

$\uparrow$ Models use dynamic per-tensor activation quantization.

4.4 Test-time compute scaling

As conventional scaling laws approach their limits, novel directions for scaling LLMs have been explored. In test-time compute scaling, model performance is improved by increasing the amount of compute used by the model at inference time. As AIMC is rather unsuitable for training, but promises orders of magnitudes higher power efficiency for inference, shifting the compute budget from training to inference is an ideal trend for AIMC. Therefore, we investigated whether our models also performed on-par to 4-bit weight quantized models when the compute used for inference is scaled up. To test this, we follow Snell et al. [74] and generate n responses to each prompt from the MATH-500 [68] dataset. Using a math process reward model [75, 76], answers are assigned a reward and the best answer is chosen by performing weighted majority voting or by simply picking the answer with the highest reward. For each model, we pick the strategy that performed best. As figure 4 shows, the analog foundation models with noisy weights perform better compared to models trained with 8-bit static input and 4-bit per-channel weight quantization. Additionally, we observe that the gap between the noisy analog foundation model and the original off-the-shelf model decreases as we increase n . While this increase is small (0.4%) for *Phi-3-mini-4k-instruct*, it is much bigger for *Llama-3.2-1B-Instruct* (3.58%).

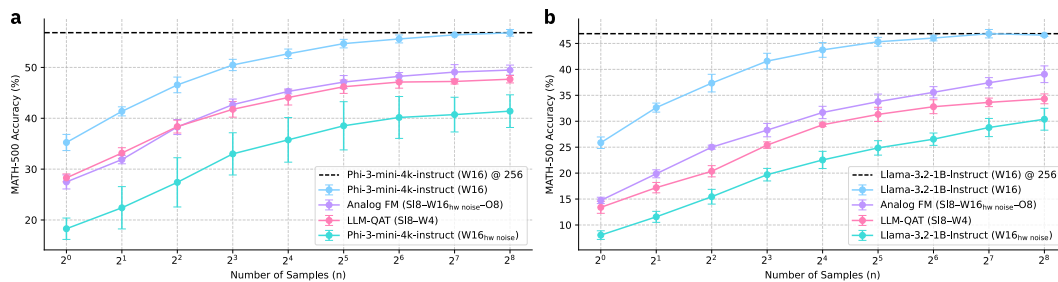


Figure 4: Math-500 accuracy (y-axis) as a function of number of generations (x-axis). Results shown for the original models, LLM-QAT models, and analog foundation models based on **a.** *Phi-3-mini-4k-instruct*- and **b.** *Llama-3.2-1B-Instruct*-based models.

5 Conclusion and Limitations

We demonstrate for the first time on a broad range of benchmarks that LLMs can be made robust against noise commonly found in AIMC hardware, achieving similar performance to models with 4-bit per-channel weight quantization and static 8-bit activation quantization. Additionally, as a byproduct of our training pipeline, analog foundation models are also robust to weight-quantization noise, enabling deployment on low-precision digital accelerators without the need for further training. Finally, we demonstrate that our analog foundation models scale better compared to their quantized counterparts when the amount of test-time compute is increased.

While our paper lays an important foundation, it also comes with limitations. Training billion-parameter models is resource intensive, especially for larger models. Although the training cost is one-time and our method only requires training on less than 1% of the total pre-training tokens, it might still be infeasible for many researchers to train their own models. Adding to this, input/output quantization and noise injection make the linear kernel approximately three times slower compared to the vanilla version, and, as we show in appendix H, training with noise injection and output quantization leads to slower convergence. As our ablation studies on the number of pre-training tokens show, the benefit of increasing the number of tokens diminishes around 20B tokens, and other avenues, such as improving data quality, must be explored to further reduce the performance gap of our models to the FP16 baseline. To reduce training overhead, one could think about employing low-rank adaptations in a HWA manner for training the network, or even better, avoid training altogether and use a post-training method for increasing the robustness of LLMs, similar to PTQ methods. Besides the high resource intensity, another limitation of our work is the relatively large gap in accuracy compared to off-the-shelf models, especially for reasoning tasks such as GSM8K or MATH-500. Further research is needed to reduce this gap. Finally, while we show that our models inherited the safety measures of the original models without significant loss in performance, it should be noted that the risk of LLMs producing toxic or harmful content still persists.

Our work answers the important question whether LLMs can be made robust enough to run on AIMC-based hardware. With this result, we hope to motivate research into further scaling AIMC chip implementations and the development of methods to further increase the robustness of LLMs to nonidealities found in AIMC hardware. We believe that our work is also relevant to the broader edge AI domain, including neuromorphic computing.

Acknowledgments and Disclosure of Funding

We are grateful to the Center for Computational Innovation at Rensselaer Polytechnic Institute for computational resources on the AiMOS Supercomputer. We also acknowledge Vijay Narayanan, Robert Haas, Jeff Burns, and Mukesh Khare for their managerial support. This research was supported by the IBM Research AI Hardware Center.

References

- [1] Jack Choquette. NVIDIA Hopper H100 GPU: Scaling Performance. *IEEE Micro*, 43(3):9–17, 2023. doi: 10.1109/MM.2023.3256796.
- [2] *AMD Instinct MI325X Accelerator*. Advanced Micro Devices, 10 2024. URL <https://www.amd.com/content/dam/amd/en/documents/instinct-tech-docs/product-briefs/instinct-mi325x-datasheet.pdf>.
- [3] Norman P. Jouppi, Cliff Young, Nishant Patil, David Patterson, Gaurav Agrawal, Raminder Bajwa, Sarah Bates, Suresh Bhatia, Nan Boden, Al Borchers, Rick Boyle, Pierre-luc Cantin, Clifford Chao, Chris Clark, Jeremy Coriell, Mike Daley, Matt Dau, Jeffrey Dean, Ben Gelb, Tara Vazir Ghaemmaghami, Rajendra Gottipati, William Gulland, Robert Hagmann, C. Richard Ho, Doug Hogberg, John Hu, Robert Hundt, Dan Hurt, Julian Ibarz, Aaron Jaffey, Alek Jaworski, Alexander Kaplan, Harshit Khaitan, Daniel Killebrew, Andy Koch, Naveen Kumar, Steve Lacy, James Laudon, James Law, Diemthu Le, Chris Leary, Zhuyuan Liu, Kyle Lucke, Alan Lundin, Gordon MacKean, Adriana Maggiore, Maire Mahony, Kieran Miller, Rahul Nagarajan, Ravi Narayanaswami, Ray Ni, Kathy Nix, Thomas Norrie, Mark Omernick, Narayana Penukonda, Andy Phelps, Jonathan Ross, Matt Ross, Amir Salek, Emad Samadiani, Chris Severn, Gregory Sizikov, Matthew Snelham, Jed Souter, Dan Steinberg, Andy Swing, Mercedes Tan, Gregory Thorson, Bo Tian, Horia Toma, Erick Tuttle, Vijay Vasudevan, Richard Walter, Walter Wang, Eric Wilcox, and Doe Hyun Yoon. In-datacenter performance analysis of a tensor processing unit. *SIGARCH Comput. Archit. News*, 45(2):1–12, June 2017. ISSN 0163-5964. doi: 10.1145/3140659.3080246. URL <https://doi.org/10.1145/3140659.3080246>.
- [4] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In I. Guyon, U. Von Luxburg, S. Bengio, H. Wallach, R. Fergus, S. Vishwanathan, and R. Garnett, editors, *Advances in Neural Information Processing Systems*, volume 30. Curran Associates, Inc., 2017. URL https://proceedings.neurips.cc/paper_files/paper/2017/file/3f5ee243547dee91fbd053c1c4a845aa-Paper.pdf.
- [5] William Fedus, Barret Zoph, and Noam Shazeer. Switch transformers: Scaling to trillion parameter models with simple and efficient sparsity. *Journal of Machine Learning Research*, 23 (120):1–39, 2022. URL <http://jmlr.org/papers/v23/21-0998.html>.
- [6] Aaron Grattafiori and (rest of authors omitted due to high number). The llama 3 herd of models, 2024. URL <https://arxiv.org/abs/2407.21783>.
- [7] Meta AI. The Llama 4 herd: The beginning of a new era of natively multimodal AI innovation. <https://ai.meta.com/blog/llama-4-multimodal-intelligence/>, 2025. Accessed: 2025-04-07.
- [8] Gemini Team and (rest of authors omitted due to high number). Gemini: A Family of Highly Capable Multimodal Models, 2023.
- [9] *Wafer-Scale Engine 3: The Largest Chip Ever Built*. Cerebras, 2024. URL <https://8968533.fs1.hubspotusercontent-na1.net/hubfs/8968533/Datasheets/WSE-3%20Datasheet.pdf>.
- [10] Dennis Abts, Garrin Kimmell, Andrew Ling, John Kim, Matt Boyd, Andrew Bitar, Sahil Parmar, Ibrahim Ahmed, Roberto DiCecco, David Han, John Thompson, Michael Bye, Jennifer Hwang, Jeremy Fowers, Peter Lillian, Ashwin Murthy, Elyas Mehtabuddin, Chetan Tekur, Thomas Sohmers, Kris Kang, Stephen Maresh, and Jonathan Ross. A software-defined tensor streaming multiprocessor for large-scale machine learning. In *Proceedings of the 49th Annual International Symposium on Computer Architecture*, ISCA ’22, page 567–580, New York, NY, USA, 2022. Association for Computing Machinery. ISBN 9781450386104. doi: 10.1145/3470496.3527405. URL <https://doi.org/10.1145/3470496.3527405>.
- [11] Andrew S. Cassidy, John V. Arthur, Filipp Akopyan, Alexander Andreopoulos, Rathinakumar Appuswamy, Pallab Datta, Michael V. Debole, Steven K. Esser, Carlos Ortega Otero, Jun Sawada, Brian Taba, Arnon Amir, Deepika Bablani, Peter J. Carlson, Myron D. Flickner,

- Rajamohan Gandhasri, Guillaume J. Garreau, Megumi Ito, Jennifer L. Klamo, Jeffrey A. Kusnitz, Nathaniel J. McClatchey, Jeffrey L. McKinstry, Yutaka Nakamura, Tapan K. Nayak, William P. Risk, Kai Schleupen, Ben Shaw, Jay Sivagnaname, Daniel F. Smith, Ignacio Terrizzano, Takanori Ueda, and Dharmendra Modha. 11.4 ibm northpole: An architecture for neural network inference with a 12nm chip. In *2024 IEEE International Solid-State Circuits Conference (ISSCC)*, volume 67, pages 214–215, 2024. doi: 10.1109/ISSCC49657.2024.10454451.
- [12] *Furiosa RNGD*. Furiosa AI, 2024. URL <https://furiosa.ai/renegead-spec>.
- [13] Geoffrey W. Burr, Robert M. Shelby, Abu Sebastian, Sangbum Kim, Seyoung Kim, Severin Sidler, Kumar Virwani, Masatoshi Ishii, Pritish Narayanan, Alessandro Fumarola, Lucas L. Sanches, Irem Boybat, Manuel Le Gallo, Kibong Moon, Jiyoo Woo, Hyunsang Hwang, and Yusuf Leblebici and. Neuromorphic computing using non-volatile memory. *Advances in Physics: X*, 2(1):89–124, 2017. doi: 10.1080/23746149.2016.1259585. URL <https://doi.org/10.1080/23746149.2016.1259585>.
- [14] Abu Sebastian, Manuel Le Gallo, Riduan Khaddam-Aljameh, and Evangelos Eleftheriou. Memory devices and applications for in-memory computing. *Nature Nanotechnology*, 15(7): 529–544, Jul 2020. doi: 10.1038/s41565-020-0655-z.
- [15] Axelera AI. The metis ai platform: A technical deep dive. *Axelera*, March 2025. URL <https://community.axelera.ai/product-updates/the-metis-ai-platform-a-technical-deepdive-125>. Accessed: 2025-05-06.
- [16] Hidehiro Fujiwara, Haruki Mori, Wei-Chang Zhao, Kinshuk Khare, Cheng-En Lee, Xiaochen Peng, Vineet Joshi, Chao-Kai Chuang, Shu-Huan Hsu, Takeshi Hashizume, Toshiaki Naganuma, Chen-Hung Tien, Yao-Yi Liu, Yen-Chien Lai, Chia-Fu Lee, Tan-Li Chou, Kerem Akarvardar, Saman Adham, Yih Wang, Yu-Der Chih, Yen-Huei Chen, Hung-Jen Liao, and Tsung-Yung Jonathan Chang. 34.4 a 3nm, 32.5tops/w, 55.0tops/mm2 and 3.78mb/mm2 fully-digital compute-in-memory macro supporting int12 $\times$ int12 with a parallel-mac architecture and foundry 6t-sram bit cell. In *2024 IEEE International Solid-State Circuits Conference (ISSCC)*, volume 67, pages 572–574, 2024. doi: 10.1109/ISSCC49657.2024.10454556.
- [17] David Schor. Iedm 2022: Did we just witness the death of sram? *WikiChip Fuse*, December 2022. URL <https://fuse.wikichip.org/news/7343/iedm-2022-did-we-just-witness-the-death-of-sram/>.
- [18] D. Fick. Analog compute-in-memory for ai edge inference. In *2022 International Electron Devices Meeting (IEDM)*, pages 21.8.1–21.8.4, 2022. doi: 10.1109/IEDM45625.2022.10019367.
- [19] Manuel Le Gallo, Riduan Khaddam-Aljameh, Milos Stanisavljevic, Athanasios Vasilopoulos, Benedikt Kersting, Martino Dazzi, Geethan Karunaratne, Matthias Brändli, Abhairaj Singh, Silvia M. Müller, Julian Büchel, Xavier Timoneda, Vinay Joshi, Malte J. Rasch, Urs Egger, Angelo Garofalo, Anastasios Petropoulos, Theodore Antonakopoulos, Kevin Brew, Samuel Choi, Injo Ok, Timothy Philip, Victor Chan, Claire Silvestre, Ishtiaq Ahsan, Nicole Saulnier, Vijay Narayanan, Pier Andrea Francese, Evangelos Eleftheriou, and Abu Sebastian. A 64-core mixed-signal in-memory compute chip based on phase-change memory for deep neural network inference. *Nature Electronics*, 6(9):680–693, Sep 2023. ISSN 2520-1131. doi: 10.1038/s41928-023-01010-1. URL <https://doi.org/10.1038/s41928-023-01010-1>.
- [20] S. Ambrogio, P. Narayanan, A. Okazaki, A. Fasoli, C. Mackin, K. Hosokawa, A. Nomura, T. Yasuda, A. Chen, A. Friz, M. Ishii, J. Luquin, Y. Kohda, N. Saulnier, K. Brew, S. Choi, I. Ok, T. Philip, V. Chan, C. Silvestre, I. Ahsan, V. Narayanan, H. Tsai, and G. W. Burr. An analog-AI chip for energy-efficient speech recognition and transcription. *Nature*, 620(7975):768–775, Aug 2023. ISSN 1476-4687. doi: 10.1038/s41586-023-06337-5. URL <https://doi.org/10.1038/s41586-023-06337-5>.
- [21] Weier Wan, Rajkumar Kubendran, Clemens Schaefer, Sukru Burc Eryilmaz, Wenqiang Zhang, Dabin Wu, Stephen Deiss, Priyanka Raina, He Qian, Bin Gao, Siddharth Joshi, Huaqiang Wu, H.-S. Philip Wong, and Gert Cauwenberghs. A compute-in-memory chip based on resistive random-access memory. *Nature*, 608(7923):504–512, Aug 2022. ISSN 1476-4687. doi: 10.1038/s41586-022-04992-8. URL <https://doi.org/10.1038/s41586-022-04992-8>.

- [22] Tai-Hao Wen, Je-Min Hung, Wei-Hsing Huang, Chuan-Jia Jhang, Yun-Chen Lo, Hung-Hsi Hsu, Zhao-En Ke, Yu-Chiao Chen, Yu-Hsiang Chin, Chin-I Su, Win-San Khwa, Chung-Chuan Lo, Ren-Shuo Liu, Chih-Cheng Hsieh, Kea-Tiong Tang, Mon-Shu Ho, Chung-Cheng Chou, Yu-Der Chih, Tsung-Yung Jonathan Chang, and Meng-Fan Chang. Fusion of memristor and digital compute-in-memory processing for energy-efficient edge computing. *Science*, 384(6693): 325–332, 2024. doi: 10.1126/science.adf5538. URL <https://www.science.org/doi/abs/10.1126/science.adf5538>.
- [23] Seungchul Jung, Hyungwoo Lee, Sungmeen Myung, Hyunsoo Kim, Seung Keun Yoon, Soon-Wan Kwon, Yongmin Ju, Minje Kim, Wooseok Yi, Shinhee Han, Baeseong Kwon, Boyoung Seo, Kilho Lee, Gwan-Hyeob Koh, Kangho Lee, Yoonjong Song, Changkyu Choi, Donhee Ham, and Sang Joon Kim. A crossbar array of magnetoresistive memory devices for in-memory computing. *Nature*, 601(7892):211–216, Jan 2022. ISSN 1476-4687. doi: 10.1038/s41586-021-04196-6. URL <https://doi.org/10.1038/s41586-021-04196-6>.
- [24] Han-Wen Hu, Wei-Chen Wang, Chung-Kuang Chen, Yung-Chun Lee, Bo-Rong Lin, Huai-Mu Wang, Yen-Po Lin, Yu-Chao Lin, Chih-Chang Hsieh, Chia-Ming Hu, Yi-Ting Lai, Han-Sung Chen, Yuan-Hao Chang, Hsiang-Pang Li, Tei-Wei Kuo, Keh-Chung Wang, Meng-Fan Chang, Chun-Hsiung Hung, and Chin-Yuan Lu. A 512gb in-memory-computing 3d-nand flash supporting similar-vector-matching operations on edge-ai devices. In *2022 IEEE International Solid-State Circuits Conference (ISSCC)*, volume 65, pages 138–140, 2022. doi: 10.1109/ISSCC42614.2022.9731775.
- [25] Po-Kai Hsu, Vaidehi Garg, Anni Lu, and Shimeng Yu. A heterogeneous platform for 3d nand-based in-memory hyperdimensional computing engine for genome sequencing applications. *IEEE Transactions on Circuits and Systems I: Regular Papers*, 71(4):1628–1637, 2024. doi: 10.1109/TCSI.2024.3362354.
- [26] Julian Büchel, Athanasios Vasilopoulos, William Andrew Simon, Irem Boybat, HsinYu Tsai, Geoffrey W. Burr, Hernan Castro, Bill Filippiak, Manuel Le Gallo, Abbas Rahimi, Vijay Narayanan, and Abu Sebastian. Efficient scaling of large language models with mixture of experts and 3d analog in-memory computing. *Nature Computational Science*, Jan 2025. ISSN 2662-8457. doi: 10.1038/s43588-024-00753-x. URL <https://doi.org/10.1038/s43588-024-00753-x>.
- [27] S. R. Nandakumar, Irem Boybat, Jin-Ping Han, Stefano Ambrogio, Praneet Adusumilli, Robert L. Bruce, Matthew BrightSky, Malte Rasch, Manuel Le Gallo, and Abu Sebastian. Precision of synaptic weights programmed in phase-change memory devices for deep learning inference. In *2020 IEEE International Electron Devices Meeting (IEDM)*, pages 29.4.1–29.4.4, 2020. doi: 10.1109/IEDM13553.2020.9371990.
- [28] S. R. Nandakumar, Irem Boybat, Vinay Joshi, Christophe Piveteau, Manuel Le Gallo, Bipin Rajendran, Abu Sebastian, and Evangelos Eleftheriou. Phase-change memory models for deep learning training and inference. In *2019 26th IEEE International Conference on Electronics, Circuits and Systems (ICECS)*, pages 727–730, 2019. doi: 10.1109/ICECS46596.2019.8964852.
- [29] I. Boybat, B. Kersting, S. Ghazi Sarwat, X. Timoneda, R. L. Bruce, M. BrightSky, M. Le Gallo, and A. Sebastian. Temperature sensitivity of analog in-memory computing using phase-change memory. In *2021 IEEE International Electron Devices Meeting (IEDM)*, pages 28.3.1–28.3.4, 2021. doi: 10.1109/IEDM19574.2021.9720519.
- [30] J. Luquin, C. Mackin, S. Ambrogio, A. Chen, F. Baldi, G. Miralles, M. J. Rasch, J. Büchel, M. Lalwani, W. Ponghiran, P. Solomon, H. Tsai, G. W. Burr, and P. Narayanan. Rapid yet accurate tile-circuit and device modeling for analog in-memory computing, 2025. URL <https://arxiv.org/abs/2506.00004>.
- [31] An Chen. A comprehensive crossbar array model with solutions for line resistance and nonlinear device characteristics. *IEEE Transactions on Electron Devices*, 60(4):1318–1326, 2013. doi: 10.1109/TED.2013.2246791.
- [32] Federico Zipoli, Daniel Krebs, and Alessandro Curioni. Structural origin of resistance drift in amorphous gete. *Phys. Rev. B*, 93:115201, Mar 2016. doi: 10.1103/PhysRevB.93.115201. URL <https://link.aps.org/doi/10.1103/PhysRevB.93.115201>.

- [33] Manuel Le Gallo, Daniel Krebs, Federico Zipoli, Martin Salinga, and Abu Sebastian. Collective structural relaxation in phase-change memory devices. *Advanced Electronic Materials*, 4(9): 1700627, 2018. doi: <https://doi.org/10.1002/aelm.201700627>. URL <https://advanced.onlinelibrary.wiley.com/doi/abs/10.1002/aelm.201700627>.
- [34] A.F. Murray and P.J. Edwards. Enhanced MLP performance and fault tolerance resulting from synaptic weight noise during training. *IEEE Transactions on Neural Networks*, 5(5):792–802, 1994. doi: 10.1109/72.317730.
- [35] Vinay Joshi, Manuel Le Gallo, Simon Haefeli, Irem Boybat, S. R. Nandakumar, Christophe Piveteau, Martino Dazzi, Bipin Rajendran, Abu Sebastian, and Evangelos Eleftheriou. Accurate deep neural network inference using computational phase-change memory. *Nature Communications*, 11(1):2473, 2020. doi: 10.1038/s41467-020-16108-9.
- [36] Julian Büchel, Fynn Firouz Faber, and Dylan Richard Muir. Network insensitivity to parameter noise via adversarial regularization. In *International Conference on Learning Representations*, 2022. URL <https://openreview.net/forum?id=-8sBpe7rDiV>.
- [37] Malte J. Rasch, Charles Mackin, Manuel Le Gallo, An Chen, Andrea Fasoli, Frédéric Odermatt, Ning Li, S. R. Nandakumar, Pritish Narayanan, Hsinyu Tsai, Geoffrey W. Burr, Abu Sebastian, and Vijay Narayanan. Hardware-aware training for large-scale and diverse deep learning inference workloads using in-memory computing-based accelerators. *Nature Communications*, 14(1):5282, 2023. doi: 10.1038/s41467-023-40770-4.
- [38] Ruihua Yu, Ze Wang, Qi Liu, Bin Gao, Zhenqi Hao, Tao Guo, Sanchuan Ding, Junyang Zhang, Qi Qin, Dong Wu, Peng Yao, Qingtian Zhang, Jianshi Tang, He Qian, and Huaqiang Wu. A full-stack memristor-based computation-in-memory system with software-hardware co-development. *Nature Communications*, 16(1):2123, Mar 2025. ISSN 2041-1723. doi: 10.1038/s41467-025-57183-0. URL <https://doi.org/10.1038/s41467-025-57183-0>.
- [39] Hang-Ting Lue, Chun-Hsiung Hung, Keh-Chung Wang, and Chih-Yuan Lu. Prospects of computing in or near flash memories. In *2024 IEEE International Electron Devices Meeting (IEDM)*, pages 1–4, 2024. doi: 10.1109/IEDM50854.2024.10873502.
- [40] Marah Abdin, Jyoti Aneja, Hany Awadalla, Ahmed Awadallah, Ammar Ahmad Awan, Nguyen Bach, Amit Bahree, Arash Bakhtiari, Jianmin Bao, Harkirat Behl, Alon Benhaim, Misha Bilenko, Johan Bjorck, Sébastien Bubeck, Martin Cai, Qin Cai, Vishrav Chaudhary, Dong Chen, Dongdong Chen, Weizhu Chen, Yen-Chun Chen, Yi-Ling Chen, Hao Cheng, Parul Chopra, Xiyang Dai, Matthew Dixon, Ronen Eldan, Victor Fragoso, Jianfeng Gao, Mei Gao, Min Gao, Amit Garg, Allie Del Giorno, Abhishek Goswami, Suriya Gunasekar, Emman Haider, Junheng Hao, Russell J. Hewett, Wenxiang Hu, Jamie Huynh, Dan Iter, Sam Ade Jacobs, Mojan Javaheripi, Xin Jin, Nikos Karampatzakis, Piero Kauffmann, Mahoud Khademi, Dongwoo Kim, Young Jin Kim, Lev Kurilenko, James R. Lee, Yin Tat Lee, Yanzhi Li, Yunsheng Li, Chen Liang, Lars Liden, Xihui Lin, Zeqi Lin, Ce Liu, Liyuan Liu, Mengchen Liu, Weishung Liu, Xiaodong Liu, Chong Luo, Piyush Madan, Ali Mahmoudzadeh, David Majercak, Matt Mazzola, Caio César Teodoro Mendes, Arindam Mitra, Hardik Modi, Anh Nguyen, Brandon Norick, Barun Patra, Daniel Perez-Becker, Thomas Portet, Reid Pryzant, Heyang Qin, Marko Radmilac, Liliang Ren, Gustavo de Rosa, Corby Rosset, Sambudha Roy, Olatunji Ruwase, Olli Saarikivi, Amin Saied, Adil Salim, Michael Santacrose, Shital Shah, Ning Shang, Hiteshi Sharma, Yelong Shen, Swadheen Shukla, Xia Song, Masahiro Tanaka, Andrea Tupini, Praneetha Vaddamanu, Chunyu Wang, Guanhua Wang, Lijuan Wang, Shuohang Wang, Xin Wang, Yu Wang, Rachel Ward, Wen Wen, Philipp Witte, Haiping Wu, Xiaoxia Wu, Michael Wyatt, Bin Xiao, Can Xu, Jiahang Xu, Weijian Xu, Jilong Xue, Sonali Yadav, Fan Yang, Jianwei Yang, Yifan Yang, Ziyi Yang, Donghan Yu, Lu Yuan, Chenruidong Zhang, Cyril Zhang, Jianwen Zhang, Li Lyna Zhang, Yi Zhang, Yue Zhang, Yunan Zhang, and Xiren Zhou. Phi-3 technical report: A highly capable language model locally on your phone, 2024. URL <https://arxiv.org/abs/2404.14219>.
- [41] I. Boybat, T. Boesch, M. Allegra, M. Baldo, J.J. Bertolini-Agnoletto, G. W. Burr, A. Buschini, A. Cabrini, E. Calvetti, C. Cappelletta, F. Conti, E. Ferro, E. Franchi Scarselli, A. Garofalo, F. Girardi, G. Islamoglu, V. P. Jonnalagadda, G. Karunaratne, C. Lammie, M. Le Gallo, C. Li, R. Massa, A. C. Ornstein, H. Pang, M. Pasotti, B. Rajendran, A. Redaelli, I. Sanli, W. A. Simon, A. Singh, S.-P. Singh, G. Urlini, A. Vasilopoulos, R. Zurla, G. Desoli, and A. Sebastian.

- Heterogeneous embedded neural processing units utilizing pcm-based analog in-memory computing. In *2024 IEEE International Electron Devices Meeting (IEDM)*, pages 1–4, 2024. doi: 10.1109/IEDM50854.2024.10873479.
- [42] Shubham Jain, Hsinyu Tsai, Ching-Tzu Chen, Ramachandran Muralidhar, Irem Boybat, Martin M. Frank, Stanisław Woźniak, Milos Stanisavljevic, Praneet Adusumilli, Pritish Narayanan, Kohji Hosokawa, Masatoshi Ishii, Arvind Kumar, Vijay Narayanan, and Geoffrey W. Burr. A Heterogeneous and Programmable Compute-In-Memory Accelerator Architecture for Analog-AI Using Dense 2-D Mesh. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, 31(1):114–127, 2023. doi: 10.1109/TVLSI.2022.3221390.
 - [43] Zechun Liu, Barlas Oguz, Changsheng Zhao, Ernie Chang, Pierre Stock, Yashar Mehdad, Yangyang Shi, Raghuraman Krishnamoorthi, and Vikas Chandra. Llm-qat: Data-free quantization aware training for large language models, 2023. URL <https://arxiv.org/abs/2305.17888>.
 - [44] Zechun Liu, Changsheng Zhao, Igor Fedorov, Bilge Soran, Dhruv Choudhary, Raghuraman Krishnamoorthi, Vikas Chandra, Yuandong Tian, and Tijmen Blankevoort. Spinquant: LLM quantization with learned rotations. In *The Thirteenth International Conference on Learning Representations*, 2025. URL <https://openreview.net/forum?id=og06DGE6FZ>.
 - [45] Bonan Zhang, Chia-Yu Chen, and Naveen Verma. Reshape and adapt for output quantization (RAOQ): Quantization-aware training for in-memory computing systems. In Ruslan Salakhutdinov, Zico Kolter, Katherine Heller, Adrian Weller, Nuria Oliver, Jonathan Scarlett, and Felix Berkenkamp, editors, *Proceedings of the 41st International Conference on Machine Learning*, volume 235 of *Proceedings of Machine Learning Research*, pages 58739–58762. PMLR, 21–27 Jul 2024. URL <https://proceedings.mlr.press/v235/zhang24i.html>.
 - [46] Amir Gholami, Sehoon Kim, Zhen Dong, Zhewei Yao, Michael W Mahoney, and Kurt Keutzer. A survey of quantization methods for efficient neural network inference. In *Low-power computer vision*, pages 291–326. Chapman and Hall/CRC, 2022.
 - [47] Boris Murmann. Mixed-signal computing for deep neural network inference. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, 29(1):3–13, 2021. doi: 10.1109/TVLSI.2020.3020286.
 - [48] Tayfun Gokmen, Malte J. Rasch, and Wilfried Haensch. The marriage of training and inference for scaled deep learning analog hardware. In *2019 IEEE International Electron Devices Meeting (IEDM)*, pages 22.3.1–22.3.4, 2019. doi: 10.1109/IEDM19573.2019.8993573.
 - [49] Sanjay Kariyappa, Hsinyu Tsai, Katie Spoon, Stefano Ambrogio, Pritish Narayanan, Charles Mackin, An Chen, Moinuddin Qureshi, and Geoffrey W. Burr. Noise-resilient dnn: Tolerating noise in pcm-based ai accelerators via noise-aware training. *IEEE Transactions on Electron Devices*, 68(9):4356–4362, 2021. doi: 10.1109/TED.2021.3089987.
 - [50] H. Tsai, S. Ambrogio, C. Mackin, P. Narayanan, R. M. Shelby, K. Rocki, A. Chen, and G. W. Burr. Inference of long-short term memory networks at software-equivalent accuracy using 2.5m analog phase change memory devices. In *2019 Symposium on VLSI Technology*, pages T82–T83, 2019. doi: 10.23919/VLSIT.2019.8776519.
 - [51] Xiaoxuan Yang, Changming Wu, Mo Li, and Yiran Chen. Tolerating noise effects in processing-in-memory systems for neural networks: A hardware–software codesign perspective. *Advanced Intelligent Systems*, 4(8):2200029, 2022. doi: <https://doi.org/10.1002/aisy.202200029>. URL <https://advanced.onlinelibrary.wiley.com/doi/abs/10.1002/aisy.202200029>.
 - [52] Kelly Spoon, Hsinyu Tsai, An Chen, Malte J. Rasch, Stefano Ambrogio, Charles Mackin, Andrea Fasoli, Alexander M. Friz, Pritish Narayanan, Milos Stanisavljevic, and Geoffrey W. Burr. Toward software-equivalent accuracy on transformer-based deep neural networks with analog memory devices. *Frontiers in Computational Neuroscience*, 15:675741, July 2021. doi: 10.3389/fncom.2021.675741.
 - [53] Yoshua Bengio, Nicholas Léonard, and Aaron Courville. Estimating or propagating gradients through stochastic neurons for conditional computation, 2013. URL <https://arxiv.org/abs/1308.3432>.

- [54] Elias Frantar, Saleh Ashkboos, Torsten Hoefer, and Dan Alistarh. GPTQ: Accurate post-training compression for generative pretrained transformers. *arXiv preprint arXiv:2210.17323*, 2022.
- [55] Ji Lin, Jiaming Tang, Haotian Tang, Shang Yang, Xingyu Dang, and Song Han. Awq: Activation-aware weight quantization for llm compression and acceleration. *CoRR*, abs/2306.00978, 2023. URL <https://doi.org/10.48550/arXiv.2306.00978>.
- [56] Saleh Ashkboos, Amirkeivan Mohtashami, Maximilian L. Croci, Bo Li, Pashmina Cameron, Martin Jaggi, Dan Alistarh, Torsten Hoefer, and James Hensman. Quarot: Outlier-free 4-bit inference in rotated LLMs. In *The Thirty-eighth Annual Conference on Neural Information Processing Systems*, 2024. URL <https://openreview.net/forum?id=dfqsW38v1X>.
- [57] Guilherme Penedo, Hynek Kydlíček, Loubna Ben allal, Anton Lozhkov, Margaret Mitchell, Colin Raffel, Leandro Von Werra, and Thomas Wolf. The fineweb datasets: Decanting the web for the finest text data at scale. In *The Thirty-eight Conference on Neural Information Processing Systems Datasets and Benchmarks Track*, 2024. URL <https://openreview.net/forum?id=n6SCkn2QaG>.
- [58] Chuteng Zhou, Prad Kadambi, Matthew Mattina, and Paul N. Whatmough. Noisy machines: Understanding noisy neural networks and enhancing robustness to analog hardware errors using distillation, 2020. URL <https://arxiv.org/abs/2001.04974>.
- [59] Julian Büchel, William Andrew Simon, Corey Lammie, Giovanni Acampa, Kaoutar El Maghraoui, Manuel Le Gallo, and Abu Sebastian. Aihwkit-lightning: A scalable hw-aware training toolkit for analog in-memory computing. In *NeurIPS 2024 Workshop Machine Learning with new Compute Paradigms*, 2024. URL <https://openreview.net/forum?id=QNdxOgGmhR>.
- [60] Adam Paszke, Sam Gross, Francisco Massa, Adam Lerer, James Bradbury, Gregory Chanan, Trevor Killeen, Zeming Lin, Natalia Gimelshein, Luca Antiga, Alban Desmaison, Andreas Köpf, Edward Yang, Zach DeVito, Martin Raison, Alykhan Tejani, Sasank Chilamkurthy, Benoit Steiner, Lu Fang, Junjie Bai, and Soumith Chintala. Pytorch: An imperative style, high-performance deep learning library, 2019. URL <https://arxiv.org/abs/1912.01703>.
- [61] Woosuk Kwon, Zhuohan Li, Siyuan Zhuang, Ying Sheng, Lianmin Zheng, Cody Hao Yu, Joseph E. Gonzalez, Hao Zhang, and Ion Stoica. Efficient memory management for large language model serving with pagedattention, 2023. URL <https://arxiv.org/abs/2309.06180>.
- [62] Peter Clark, Isaac Cowhey, Oren Etzioni, Tushar Khot, Ashish Sabharwal, Carissa Schoenick, and Oyvind Tafjord. Think you have solved question answering? try arc, the ai2 reasoning challenge, 2018. URL <https://arxiv.org/abs/1803.05457>.
- [63] Dan Hendrycks, Collin Burns, Steven Basart, Andy Zou, Mantas Mazeika, Dawn Song, and Jacob Steinhardt. Measuring massive multitask language understanding, 2021. URL <https://arxiv.org/abs/2009.03300>.
- [64] Di Jin, Eileen Pan, Nassim Oufattole, Wei-Hung Weng, Hanyi Fang, and Peter Szolovits. What disease does this patient have? a large-scale open domain question answering dataset from medical exams, 2020. URL <https://arxiv.org/abs/2009.13081>.
- [65] Wanjun Zhong, Ruixiang Cui, Yiduo Guo, Yaobo Liang, Shuai Lu, Yanlin Wang, Amin Saied, Weizhu Chen, and Nan Duan. AGIEval: A human-centric benchmark for evaluating foundation models. In Kevin Duh, Helena Gomez, and Steven Bethard, editors, *Findings of the Association for Computational Linguistics: NAACL 2024*, pages 2299–2314, Mexico City, Mexico, June 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.findings-naacl.149. URL <https://aclanthology.org/2024.findings-naacl.149/>.
- [66] Rowan Zellers, Ari Holtzman, Yonatan Bisk, Ali Farhadi, and Yejin Choi. Hellaswag: Can a machine really finish your sentence?, 2019. URL <https://arxiv.org/abs/1905.07830>.
- [67] Karl Cobbe, Vineet Kosaraju, Mohammad Bavarian, Mark Chen, Heewoo Jun, Lukasz Kaiser, Matthias Plappert, Jerry Tworek, Jacob Hilton, Reiichiro Nakano, Christopher Hesse, and John Schulman. Training verifiers to solve math word problems, 2021. URL <https://arxiv.org/abs/2110.14168>.

- [68] Hunter Lightman, Vineet Kosaraju, Yura Burda, Harri Edwards, Bowen Baker, Teddy Lee, Jan Leike, John Schulman, Ilya Sutskever, and Karl Cobbe. Let’s verify step by step, 2023. URL <https://arxiv.org/abs/2305.20050>.
- [69] Christopher Clark, Kenton Lee, Ming-Wei Chang, Tom Kwiatkowski, Michael Collins, and Kristina Toutanova. Boolq: Exploring the surprising difficulty of natural yes/no questions, 2019. URL <https://arxiv.org/abs/1905.10044>.
- [70] Yixin Nie, Adina Williams, Emily Dinan, Mohit Bansal, Jason Weston, and Douwe Kiela. Adversarial NLI: A new benchmark for natural language understanding. In Dan Jurafsky, Joyce Chai, Natalie Schluter, and Joel Tetreault, editors, *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 4885–4901, Online, July 2020. Association for Computational Linguistics. doi: 10.18653/v1/2020.acl-main.441. URL <https://aclanthology.org/2020.acl-main.441/>.
- [71] Jeffrey Zhou, Tianjian Lu, Swaroop Mishra, Siddhartha Brahma, Sujoy Basu, Yi Luan, Denny Zhou, and Le Hou. Instruction-following evaluation for large language models, 2023. URL <https://arxiv.org/abs/2311.07911>.
- [72] Paul Röttger, Hannah Kirk, Bertie Vidgen, Giuseppe Attanasio, Federico Bianchi, and Dirk Hovy. XSTest: A test suite for identifying exaggerated safety behaviours in large language models. In Kevin Duh, Helena Gomez, and Steven Bethard, editors, *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*, pages 5377–5400, Mexico City, Mexico, June 2024. Association for Computational Linguistics. doi: 10.18653/v1/2024.naacl-long.301. URL <https://aclanthology.org/2024.naacl-long.301/>.
- [73] Manuel Le Gallo, Corey Lammie, Julian Büchel, Fabio Carta, Omobayode Fagbohunge, Charles Mackin, Hsinyu Tsai, Vijay Narayanan, Abu Sebastian, Kaoutar El Maghraoui, and Malte J. Rasch. Using the ibm analog in-memory hardware acceleration kit for neural network training and inference. *APL Machine Learning*, 1(4):041102, 11 2023. ISSN 2770-9019. doi: 10.1063/5.0168089. URL <https://doi.org/10.1063/5.0168089>.
- [74] Charlie Snell, Jaehoon Lee, Kelvin Xu, and Aviral Kumar. Scaling llm test-time compute optimally can be more effective than scaling model parameters, 2024. URL <https://arxiv.org/abs/2408.03314>.
- [75] Peiyi Wang, Lei Li, Zhihong Shao, R. X. Xu, Damai Dai, Yifei Li, Deli Chen, Y. Wu, and Zhifang Sui. Math-shepherd: Verify and reinforce llms step-by-step without human annotations, 2024. URL <https://arxiv.org/abs/2312.08935>.
- [76] Wei Xiong, Hanning Zhang, Nan Jiang, and Tong Zhang. An implementation of generative prm. <https://github.com/RLHFlow/RLHF-Reward-Modeling>, 2024.
- [77] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta: A robustly optimized bert pretraining approach, 2019. URL <https://arxiv.org/abs/1907.11692>.
- [78] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert: Pre-training of deep bidirectional transformers for language understanding, 2019. URL <https://arxiv.org/abs/1810.04805>.
- [79] Alex Wang, Amanpreet Singh, Julian Michael, Felix Hill, Omer Levy, and Samuel Bowman. GLUE: A multi-task benchmark and analysis platform for natural language understanding. In Tal Linzen, Grzegorz Chrupała, and Afra Alishahi, editors, *Proceedings of the 2018 EMNLP Workshop BlackboxNLP: Analyzing and Interpreting Neural Networks for NLP*, pages 353–355, Brussels, Belgium, November 2018. Association for Computational Linguistics. doi: 10.18653/v1/W18-5446. URL <https://aclanthology.org/W18-5446/>.
- [80] Moran Shkolnik, Brian Chmiel, Ron Banner, Gil Shomron, Yury Nahshan, Alex Bronstein, and Uri Weiser. Robust quantization: One model to rule them all, 2020. URL <https://arxiv.org/abs/2002.07686>.

- [81] Diederik P. Kingma and Jimmy Ba. Adam: A method for stochastic optimization, 2017. URL <https://arxiv.org/abs/1412.6980>.
- [82] Edward Beeching, Lewis Tunstall, and Sasha Rush. Scaling test-time compute with open models. URL <https://huggingface.co/spaces/HuggingFaceH4/blogpost-scaling-test-time-compute>.

Appendix

Contents

A	Analog RoBERTa	20
B	Ablations on the training methodology	20
B.1	Data generation	20
B.2	Number of training tokens	21
B.3	Source of training data	21
B.4	Importance of knowledge distillation	22
C	Ablations on hardware aware training	22
C.1	Output quantization	22
C.2	Noise injection	22
C.3	Weight clipping	24
D	The effect of weight clipping on signal-to-noise ratio	25
D.1	Weight clipping reduces quantization error	25
E	Evaluation details	26
E.1	Noise models used	26
E.2	Benchmarks	26
E.3	Evaluation prompting	28
F	Additional noise experiments	37
F.1	Sensitivity to noise magnitude and generalization across memory technologies . . .	37
F.2	Robustness to read noise and conductance drift	37
F.3	Input quantization precision	38
G	Impact of tiled AIMC architectures	39
H	Training convergence	39
I	Training details	40
J	Test-time compute scaling details	40
K	Hardware throughput and energy efficiency estimation	42

Table 4: Summary of each GLUE task, including the number of training and test samples, the metric used for evaluation, and the task type.

Task	Train Samples	Test Samples	Metric	Type
CoLA	8.55k	1.06k	Matthew's Correlation	Binary Classification
MNLI	393k	9.80k	Accuracy	Multi-class Classification
MRPC	3.67k	1.73k	Accuracy	Binary Classification
QNLI	105k	5.46k	Accuracy	Binary Classification
QQP	364k	40.4k	Accuracy	Binary Classification
RTE	2.49k	3.00k	Accuracy	Binary Classification
SST2	67.3k	1.82k	Accuracy	Binary Classification
STSBB	5.75k	1.38k	Pearson Correlation	Regression

Table 5: Performance comparison on GLUE benchmark of RoBERTa when HWA training is either applied during the pre-training and finetuning stage, or only during the finetuning stage.

Model	CoLA	MNLI	MRPC	QNLI	QQP	RTE	SST2	STSBB	Avg.
RoBERTa (W16)	62.19	87.82	89.46	92.81	91.72	78.34	94.72	90.92	85.98
Pre-train + finetune (S18-W16 _{hw noise})	62.48 [†] ± 1.03	87.21 ± 0.07	88.33 [†] ± 1.17	92.41 ± 0.17	91.42 ± 0.11	74.87 [†] ± 0.51	94.70 ± 0.35	90.52 ± 0.18	85.24
Finetune only (S18-W16 _{hw noise})	51.71 ± 2.01	86.53 ± 0.16	87.01 ± 0.51	92.12 ± 0.21	91.29 ± 0.05	73.61 ± 2.01	93.89 ± 0.35	90.04 ± 0.15	83.27

[†] No noise injection during training was used in the finetuning stage.

A Analog RoBERTa

In this section, we show that employing HWA training only during the finetuning stage yields inferior results compared to when HWA training is already performed during the pre-training stage, especially for tasks where the amount of finetuning data is scarce. To show this, we use RoBERTa [77], a variant of BERT [78] that improves on BERT's performance by training on 10× more data. To compare the performance when HWA training is employed during the pre-training stage or only during the finetuning stage, we train an analog version of RoBERTa by repeating the pre-training process on roughly 20% of the pre-training tokens. More specifically, we use a mix of BookCorpus, English Wikipedia, CC-News, and OpenWebText to train the model for 31'000 steps on 8 V100 GPUs, which takes about 35 hours. Using this model as the starting point, we then finetune it on each GLUE [79] (see table 4 for details) task starting from five different seeds.

During evaluation, we use the same hardware-realistic noise used throughout this paper and evaluate each model on five seeds. More details can be found in appendix E. Table 5 shows that by applying HWA training during the pre-training stage, performance can be improved by 1.97% on average, bringing the performance under hardware-realistic noise within 1% of the FP16 performance of the original RoBERTa model. Significant gains in performance can be observed for tasks with fewer training samples such as CoLA, MRPC, or RTE. From this, we conclude that it is generally better to apply HWA training in the pre-training phase of a model.

B Ablations on the training methodology

In this section, we show the results of our ablation studies on the training methodology we used. We investigate different data generation methods, the choice between real world and synthetic data, different loss functions, and the amount of tokens used for training.

B.1 Data generation

For all models, sequences of length 4096 are sampled. We continue sampling even when an EOS token was encountered. When sampling from the softmax distribution, we sample from the top-50 values in *Llama-3.2-1B-Instruct* and from all of the values in *Phi-3-mini-4k-instruct*. Similar to Liu et al. [43], we compare three different data generation strategies. The first strategy is based on sampling every token from the softmax distribution, and is denoted as "SSS". The second strategy – denoted "RGS" – first samples the initial token uniformly at random, followed by sampling the next 5 tokens greedily, and finally sampling the rest of the tokens from the softmax distribution. The last strategy – denoted "SGS" – initially samples the first token from the softmax distribution, followed by greedily picking the next 5 tokens, followed by sampling the rest of the tokens from the softmax. Table 6 shows benchmark results for *Phi-3-mini-4k-instruct* analog foundation models

trained without output quantization on 1B tokens. We found that while the differences between the methods are small, pure softmax sampling performed the best.

Table 6: Results on varying synthetic data generation techniques on *Phi-3-mini-4k-instruct*.

Model	Generation Technique	Benchmarks								Avg.
		MMLU (5-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
SI8-W16	SSS	65.37	78.11	80.56	47.03	36.15	82.20	89.87	50.81	66.26
SI8-W16	RGS	65.15	76.22	78.80	47.03	35.70	82.54	90.00	49.81	65.66
SI8-W16	SGS	65.64	75.95	80.04	45.63	36.25	82.29	90.13	49.44	65.67

Upon further inspection, we found that some generated sequences collapsed on repetition of random characters or sequences. We therefore also experimented with filtering out the 20% sequences with the lowest log-probability. While this boosted performance by 0.45%, we did not use it in our final models as the difference was too insignificant and data generation was slower. We believe that more experimentation in this direction could yield improvements of the order of 1%.

B.2 Number of training tokens

Using the optimal HWA training and data generation configuration, we trained both analog foundation models on a varying number of synthetically generated tokens. For *Llama-3.2-1B-Instruct* we trained on up to 40B tokens because training was almost twice as fast as for *Phi-3-mini-4k-instruct*, for which we trained on up to 20B tokens. Table 7 shows the benchmark results for the varying number of tokens. Overall, the best performance is reached for 20B tokens. Interestingly, for *Llama-3.2-1B-Instruct* we see that the performance with 40B tokens declines on average compared to the 20B case. We observe that on the reasoning tasks like GSM8K, performance still improved, but declined on tasks based on factual knowledge like BoolQ.

Table 7: Ablation study on the effect of number of training tokens for *Phi-3-mini-4k-instruct* and *Llama-3.2-1B-Instruct*.

Model	Tokens	Benchmarks									Avg.
		MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
<i>Phi-3-mini-4k-instruct</i>											
Analog FM (SI8-W16-O8)	1B	65.50	70.74	76.73	79.90	47.48	36.56	82.25	90.40	51.09	66.74
Analog FM (SI8-W16 _{low noise} -O8)	1B	63.18±0.34	63.62±1.49	74.71±1.82	76.31±1.20	43.49±0.80	34.21±0.89	81.34±0.63	89.40±0.34	46.80±1.93	63.67
Analog FM (SI8-W16-O8)	10B	66.79	76.27	76.54	82.35	48.19	36.85	83.36	90.66	49.97	67.89
Analog FM (SI8-W16 _{low noise} -O8)	10B	65.02±0.31	71.79±0.95	75.22±1.84	79.89±0.92	46.57±1.13	35.55±0.59	82.38±0.57	89.77±0.35	47.43±2.38	65.96
Analog FM (SI8-W16-O8)	20B	67.24	77.41	77.00	83.05	49.76	37.20	83.62	90.74	51.94	68.66
Analog FM (SI8-W16 _{low noise} -O8)	20B	65.14±0.33	71.63±0.84	75.44±1.71	80.25±0.68	46.54±0.81	35.87±0.62	82.75±0.59	89.93±0.26	49.38±1.93	66.33
<i>Llama-3.2-1B-Instruct</i>											
Analog FM (SI8-W16-O8)	1B	45.16	33.74	66.12	31.38	31.76	26.05	46.67	65.03	35.34	42.36
Analog FM (SI8-W16 _{low noise} -O8)	1B	40.83 ± 0.31	25.59 ± 1.38	63.28 ± 0.47	28.35 ± 0.35	28.11 ± 1.09	25.16 ± 0.36	44.17 ± 0.93	59.19 ± 0.75	33.79 ± 1.10	38.72
Analog FM (SI8-W16-O8)	10B	44.92	35.94	64.65	30.80	33.33	25.92	50.34	65.82	36.47	43.13
Analog FM (SI8-W16 _{low noise} -O8)	10B	42.56 ± 0.27	29.64 ± 0.96	63.40 ± 0.48	29.06 ± 0.37	29.94 ± 1.21	25.80 ± 0.38	47.01 ± 1.01	62.31 ± 0.41	34.48 ± 1.19	40.47
Analog FM (SI8-W16-O8)	20B	44.94	36.32	64.92	30.75	32.55	26.50	50.09	67.21	35.72	43.22
Analog FM (SI8-W16 _{low noise} -O8)	20B	42.83 ± 1.08	30.46 ± 1.30	62.77 ± 0.91	29.34 ± 1.09	31.21 ± 1.29	26.26 ± 0.72	47.34 ± 1.77	63.04 ± 2.00	34.92 ± 0.96	40.91
Analog FM (SI8-W16-O8)	40B	45.10	37.83	61.90	28.35	32.08	26.24	51.28	67.21	36.25	42.92
Analog FM (SI8-W16 _{low noise} -O8)	40B	43.21 ± 0.75	31.29 ± 1.19	61.06 ± 1.14	27.36 ± 0.69	31.08 ± 1.65	25.63 ± 0.96	47.70 ± 1.62	63.74 ± 1.07	35.12 ± 0.85	40.69

For standard QAT, a similar performance trend can be observed as table 8 demonstrates.

Table 8: Ablation study on the number of tokens for LLM-QAT.

Model	Tokens	Benchmarks								Avg.	
		MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)		ANLI (7-shot)
LLM-QAT (SI8-W4)	1B	64.16	67.17	76.94	78.59	46.31	35.47	81.06	88.80	48.84	65.26
LLM-QAT (SI8-W4)	10B	65.27	71.95	74.50	81.09	45.75	35.23	82.34	89.18	50.25	66.17
LLM-QAT (SI8-W4)	20B	64.12	68.92	75.11	79.22	45.20	36.30	81.48	89.18	51.16	65.63

B.3 Source of training data

In order to determine whether it is actually beneficial to use synthetic data, we trained two analog foundation models, one on a 1B token subset of the FineWeb [57] dataset and the other on 1B

synthetically generated tokens. As table 9 shows, training on synthetically generated data leads to better performance, however, when resources to generate the synthetic dataset are not available, using a publicly available dataset of high quality still leads to decent results.

Table 9: Ablation study on the choice of the data source on *Phi-3-mini-4k-instruct*.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
SI8-W16, FineWeb	66.89	68.31	81.28	75.63	49.14	36.04	83.62	90.53	49.34	66.75
SI8-W16, Synthetic	66.22	71.46	77.78	80.90	49.53	35.65	83.33	90.57	51.25	67.41

B.4 Importance of knowledge distillation

By employing harsh HWA training or QAT, information and capability stored in a pre-trained LLM is lost and needs to be recovered during the re-training process. When using standard cross-entropy during this re-training process, the LLM will start to model the data used for this stage, which is undesirable as the model will represent a different data distribution compared to the original model which is almost always trained on data that is not publicly available. By training on a pure distillation loss, the model is only encouraged to imitate the original model on the given data – and not to model it. Table 10 shows the results when *Phi-3-mini-4k-instruct* is trained with and without distillation on a 1B subset of the FineWeb dataset. Training without distillation leads to an average drop of 8.05%.

Table 10: Ablation study on the choice of loss function for re-training on *Phi-3-mini-4k-instruct*.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
SI8-W16, Distillation	67.09	69.90	80.09	75.39	51.02	36.35	83.70	90.53	48.69	66.97
SI8-W16, No distillation	62.28	55.27	71.22	56.86	43.24	31.49	80.55	87.92	41.47	58.92

C Ablations on hardware aware training

C.1 Output quantization

To quantify the drop in performance attributed to globally static output quantization, we trained two *Phi-3-mini-4k-instruct*-based models on 1B synthetically generated tokens with and without output quantization. Table 11 shows the results across the 9 benchmarks. As can be seen, adding 8-bit globally static output quantization during training leads to a degradation of 0.36% when no noise is applied to the weights at inference time, and 0.27% when noise is applied.

Table 11: The effect of output quantization on the performance of *Phi-3-mini-4k-instruct*.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
SI8-W16	67.69	76.12	77.89	82.67	49.84	37.08	84.47	90.91	50.03	68.52
SI8-W16 _{noisy}	65.23	70.66	75.82	78.97	46.16	35.75	82.30	89.81	46.66	65.71
SI8-W16-O8	67.05	75.89	77.31	81.70	49.14	37.62	84.56	90.87	49.28	68.16
SI8-W16 _{noisy} -O8	64.73	69.90	76.27	77.94	45.50	35.74	82.27	89.92	46.67	65.44

C.2 Noise injection

An important aspect of HWA training is noise injection into the weights during the forward pass in order to make the model more robust to small weight fluctuations. Adding noise during training of LLMs generally lowers downstream FP16 performance. In contrast to earlier works that have claimed better generalization as a result of noise injection, we have not observed improvement of FP16 performance. We have, however, observed an improvement in performance when noise is added during evaluation. This leads to a trade-off between FP16 accuracy and robustness: adding too much

noise will degrade the FP16 performance too much, while adding no noise will yield a large drop in accuracy when noise is added during evaluation. To study the optimal noise injection magnitude, we sweep the amount of noise injected during training of various *Phi-3-mini-4k-instruct*-based models, which are trained with weight clipping ($\alpha = 3.0$) and 8-bit input- and output quantization. This amount is controlled by the parameter γ_{weight} in eq. 3, which we re-state here.

$$\mathbf{W}_{:,i}^{\text{noisy}} \leftarrow \mathbf{W}_{:,i} + \gamma_{\text{weight}} \cdot \max(|\mathbf{W}_{:,i}|) \cdot \tau$$

where $\tau \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

Figure 5 shows the average accuracy of the benchmarks (y-axis) for models trained with different amounts of noise injection (x-axis). As the amount of training noise increases, the gap between the noise-free accuracy (blue) and noisy accuracy (pink, 10 seeds, only mean is shown) shrinks, which is a sign of increased robustness. This, however, comes at the cost of lower FP16 performance, clearly illustrating the trade-off between accuracy and robustness. It can be seen that the optimal robustness is achieved for $\gamma_{\text{weight}} = 0.02$, which was also used in the final analog foundation model training run for *Phi-3-mini-4k-instruct*.

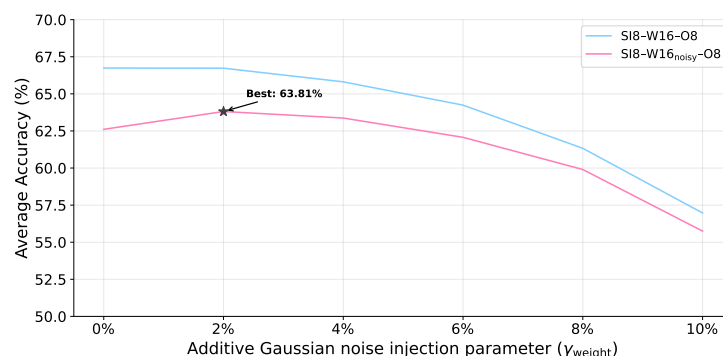


Figure 5: Sweep over amount of noise injected during training.

Besides the amount of noise injected into the model, we found that the type of noise also matters. We performed further studies on three different types of noise: purely additive Gaussian noise, purely multiplicative Gaussian noise, and a combination of the two, which we name affine. These types can be captured in eq. 5 with three different hyperparameter configurations. For additive Gaussian noise, we have $\beta_{\text{weight}} = 0$, for multiplicative noise we have $\gamma_{\text{weight}} = 0$, and for the affine type of noise we have $\beta_{\text{weight}} \neq 0$ and $\gamma_{\text{weight}} \neq 0$.

$$\mathbf{W}_{:,i}^{\text{noisy}} \leftarrow \mathbf{W}_{:,i} + (\gamma_{\text{weight}} \cdot \max(|\mathbf{W}_{:,i}|) + \beta_{\text{weight}} |\mathbf{W}_{:,i}|) \cdot \tau$$

where $\tau \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$

(5)

Using these noise types, we performed the following sweeps on *Phi-3-mini-4k-instruct* using 1B tokens for training. For the additive noise we swept γ_{weight} , for the multiplicative noise we swept β_{weight} , and for the affine noise we kept $\gamma_{\text{weight}} = 0.02$ and swept β_{weight} . We generally found that the multiplicative component did not contribute any robustness, which was surprising as our assumed noise model is most closely modeled by the affine, and not by the additive, noise type. This result suggests that it is more important to increase the robustness of the small weights, as they receive the most relative noise during evaluation, and that the larger weights are already robust enough. Having found the optimal values for the additive and affine noise type, we trained two models on 10B tokens to test which of the two types yield the best results. Table 12 shows that both noise configurations perform similarly well, especially on hard tasks such as GSM8K, compared to the baseline that was trained with no noise injection.

We conclude that the additive component of the noise model is the most important, and that training with heteroscedastic noise does not offer any advantage in our setting. For this reason, and because constant additive noise is slightly more computationally efficient, we used constant additive noise for the training of our analog foundation models.

Table 12: Performance comparisons across different noise injection settings during HWA training on *Phi-3-4096-Instruct*. All evaluation with noise were performed across 10 runs with different seeds, and the reported values represent the mean performance. The models were trained with 10 billion tokens. We compare three different models: "No noise" ($\beta_{\text{weight}} = \gamma_{\text{weight}} = 0\%$), "Affine" ($\beta_{\text{weight}} = 6\%$, $\gamma_{\text{weight}} = 2\%$), and "Additive" ($\beta_{\text{weight}} = 0\%$, $\gamma_{\text{weight}} = 2\%$).

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
No Noise (SI8-W16-O8)	67.05	75.89	77.31	81.70	49.14	37.62	84.56	90.87	49.28	68.16
No Noise (SI8-W16 _{noisy} -O8)	64.73	69.90	76.27	77.94	45.50	35.74	82.27	89.92	46.67	65.44
Affine (SI8-W16-O8)	66.56	75.74	77.16	82.16	48.03	36.41	83.62	90.40	51.22	67.93
Affine (SI8-W16 _{noisy} -O8)	64.72	71.74	75.83	80.20	46.00	35.27	82.34	89.79	47.69	65.95
Constant (SI8-W16-O8)	66.79	76.57	76.27	82.35	49.29	36.72	83.28	90.61	49.88	67.97
Constant (SI8-W16 _{noisy} -O8)	65.02	72.00	75.28	79.76	46.36	35.63	82.35	89.80	47.49	65.96

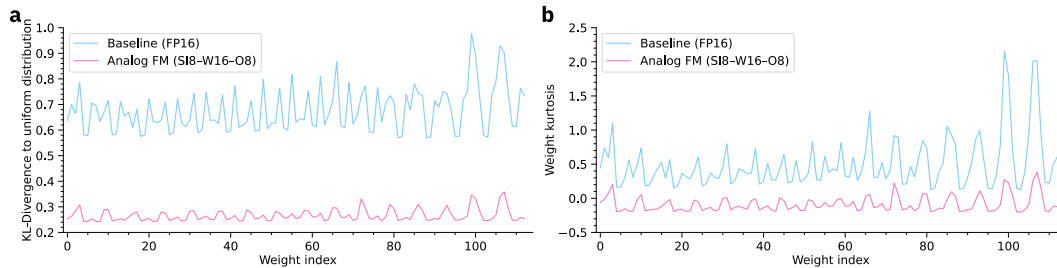


Figure 6: Analog foundation models trained with iterative weight clipping have a much smaller KL-divergence to the uniform distribution compared to the baseline models (a.). This is achieved by iteratively removing outliers which also reduces Kurtosis (b.), another proxy for the KL-divergence often used to increase robustness to quantization noise [80].

C.3 Weight clipping

During training, weights are clamped to $\pm\alpha$ standard deviations. By clamping after every optimizer step, outliers are iteratively removed and weights are prevented from becoming outliers. The hyperparameter α acts as a regularization parameter that controls how tight the weight distribution should be. We generally find that values between 2.0 and 3.5 work best, as they remove outliers without impacting model performance much. In Shkolnik et al. [80], the authors have shown that uniform distributions are more robust to changes in the quantization step size compared to normal distributions. Because the authors showed that this assumption holds for all quantization step sizes, we hypothesize that there is also a connection to our noise models. Building on their insight, the authors propose Kurtosis regularization as a proxy for bringing the weight distributions closer to the uniform distribution. Our method is an equally valid and computationally even cheaper proxy to bringing the weight distribution closer to the uniform distribution as figure 6 shows.

Table 13 shows the performance for models based on *Phi-3-4096-Instruct* trained on 10B synthetically generated tokens with only clipping and with clipping and noise injection. As can be seen, the average improvement from adding clipping to the training is 2.52% while the improvement from adding noise injection is only 0.52%.

Table 13: Comparison between the effect of clipping and the effect of noise injection during training of *Phi-3-4096-Instruct*-based models on 10B tokens. For noisy evaluations, experiments were repeated for 10 seeds and only the mean is shown.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
Baseline (W16)	69.36	79.91	78.62	83.51	52.91	38.06	84.56	91.08	52.25	70.03
<i>Phi-3-4096-Instruct</i> (W16 _{hw noise})	64.06	60.24	73.72	75.72	45.66	33.94	81.24	89.34	42.87	62.92
Clipping (SI8-W16-O8)	67.05	75.89	77.31	81.70	49.14	37.62	84.56	90.87	49.28	68.16
Clipping (SI8-W16 _{noisy} -O8)	64.73	69.90	76.27	77.94	45.50	35.74	82.27	89.92	46.67	65.44 (+2.52%)
Clipping + Noise (SI8-W16-O8)	66.79	76.57	76.27	82.35	49.29	36.72	83.28	90.61	49.88	67.97
Clipping + Noise (SI8-W16 _{noisy} -O8)	65.02	72.00	75.28	79.76	46.36	35.63	82.35	89.80	47.49	65.96 (+0.52%)

D The effect of weight clipping on signal-to-noise ratio

Weight clipping applied during training has an important effect on the distribution of weights and their robustness to noise. By iteratively clamping weights to $\pm\alpha$ standard deviations after each optimizer step (see eq. 4), the weight distribution becomes tighter and outliers are removed. This has a beneficial side effect: smaller weight magnitudes are mapped to relatively larger conductance values on the device, which improves their signal-to-noise ratio.

Figure 7a shows the conductance-dependent SNR characteristic of the PCM noise model used in our experiments. The SNR is highest for mid-range conductance values and degrades for both very small and very large conductances, with particularly poor SNR near zero conductance. By tightening the weight distribution through clipping, our analog foundation models shift the weight-to-conductance mapping such that more weights occupy the higher-SNR regions of the device characteristic.

We quantified this effect by computing the average per-layer mean SNR for each model, where the SNR of each weight is determined by its normalized conductance value according to the PCM noise model. Figures 7b and 7c show the SNR distribution across all linear layers for *Llama-3.2-1B-Instruct* and *Phi-3-mini-4k-instruct*, respectively. The analog foundation models (purple) consistently achieve higher average SNR compared to both the off-the-shelf models (cyan) and the LLM-QAT models (pink). Specifically, for *Llama-3.2-1B-Instruct*, the analog foundation model achieves an average per-layer mean SNR of 14.46 dB compared to 11.66 dB for the off-the-shelf model and 12.02 dB for the LLM-QAT model. For *Phi-3-mini-4k-instruct*, the analog foundation model achieves 13.66 dB compared to 12.02 dB and 12.18 dB, respectively.

Interestingly, the LLM-QAT models also show improved SNR compared to the off-the-shelf models, which partially explains their observed robustness to analog noise. However, their SNR improvement is more modest than that of the analog foundation models.

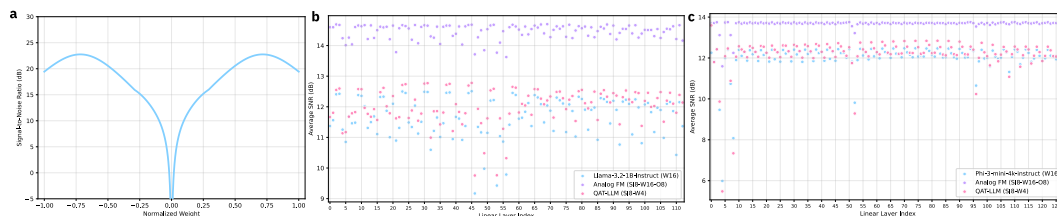


Figure 7: **a.** Signal-to-noise ratio as a function of normalized conductance for the PCM noise model. The SNR is highest for mid-range conductances and degrades near zero and at the extremes. **b.** Average per-layer SNR for each linear layer in *Llama-3.2-1B-Instruct*-based models. **c.** Average per-layer SNR for each linear layer in *Phi-3-mini-4k-instruct*-based models. The analog foundation models (purple) consistently achieve higher SNR compared to off-the-shelf models (cyan) and LLM-QAT models (pink).

D.1 Weight clipping reduces quantization error

While noise injection during training encourages weights to converge to flatter regions of the loss landscape [35], we find that iterative weight clipping (see eq. 4) contributes most significantly to robustness against weight quantization. Weight clipping explicitly removes outliers from the weight distribution, which reduces the dynamic range that must be covered by the quantization grid. This leads to smaller quantization errors compared to models where outliers remain in the weight distribution.

To quantify this effect, we computed the mean absolute quantization error for each linear layer when applying 4-bit per-channel RTN quantization to the weights. Figure 8 shows the per-layer quantization errors for both *Phi-3-mini-4k-instruct* (figure 8a) and *Llama-3.2-1B-Instruct* (figure 8b). The analog foundation models (purple) consistently exhibit lower quantization errors compared to the LLM-QAT models (pink) across nearly all layers. For *Phi-3-mini-4k-instruct*, the analog foundation model achieves an average per-layer mean absolute quantization error of 0.0034 ± 0.0005 compared to 0.0062 ± 0.0013 for the LLM-QAT model. For *Llama-3.2-1B-Instruct*, the corresponding values are 0.0019 ± 0.0006 and 0.0045 ± 0.0015 , respectively.

This substantial reduction in quantization error—approximately 45% for *Phi-3-mini-4k-instruct* and 58% for *Llama-3.2-1B-Instruct*—helps explain why our analog foundation models perform competitively when deployed on low-precision digital hardware (see table 3). The tighter weight distributions resulting from iterative clipping allow the quantization grid to more efficiently represent the weight values, leading to smaller approximation errors.

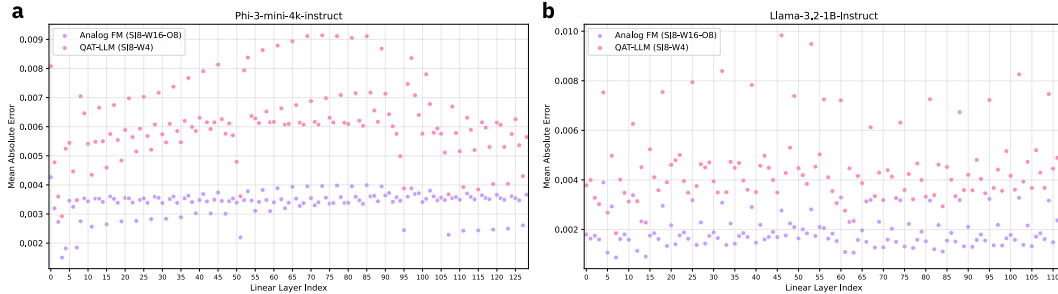


Figure 8: Mean absolute quantization error per linear layer when applying 4-bit per-channel round-to-nearest quantization. **a.** Results for *Phi-3-mini-4k-instruct*-based models. **b.** Results for *Llama-3.2-1B-Instruct*-based models. Analog foundation models (purple) consistently achieve lower quantization errors compared to LLM-QAT models (pink) due to tighter weight distributions resulting from iterative clipping during training.

E Evaluation details

E.1 Noise models used

We use the noise model from the IBM Hermes Project chip [19], a 64-core PCM-based state-of-the-art AIMC chip. In this chip, weights are represented in "unit cells" comprising of 4 PCM devices per unit cell, with two devices per polarity, meaning that two devices can be used to encode one weight. When two devices are used instead of one, programming noise is reduced by roughly a factor of $\sqrt{2}$. In this paper, we assume the case where two devices are used to encode a weight. Figure 9 shows the weight error $\text{std}(W - \hat{W})/W_{\max}$ as a function of the normalized weight, where $\text{std}(\cdot)$ is the standard deviation over all elements in the error matrix, W is the target weight matrix, $\hat{W}$ is the inferred programmed weight matrix, and $W_{\max}$ is the maximum weight. The error bars in the plot indicate one standard deviation over measurements of all 64 cores of the chip. The model we used is a third-degree polynomial fitted to this data:

$$\begin{aligned} W_{\text{hw noise } i,j} &= W_{i,j} + \eta \\ \eta &\sim \mathcal{N}(0, \sigma_{i,j}^2) \\ \sigma_{i,j} &= 1.23e - 5W_{i,j}^3 - 3.06e - 3W_{i,j}^2 + 2.45e - 1W_{i,j} + 2.11 \end{aligned}$$

E.2 Benchmarks

More details about the used benchmarks are outlined in table 14. In the following, we show the exact shots used for each evaluation, detail how the data is pre-processed, and how the prediction is extracted from the model response.

Table 14: Benchmarks used for evaluation, showing number of test samples, number of few-shot examples used, and task types. MC denotes multiple-choice questions, and CoT refers to chain-of-thought reasoning.

	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	HellaSwag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	MATH-500 (0-shot)	IFEval (0-shot)	XSTest (0-shot)
Test Samples	6,168	3,200	1,176	2,376	6,552	1,328	10,056	2,544	14,064	500	541	450
Type	4 MC	CoT	Yes/No	4 MC	5 MC	4 MC	4 MC	4 MC	3 MC	CoT	Prompt/instruct	VPRR/IPRR
Chance (%)	26.89	0.0*	62.17	25.73	21.93	25.24	26.71	26.60	0.0*	0.0*	0.0*	0.0 [†]

[†] Assuming if no reasonable answer is generated the response is classified as "3_partial_refusal", which is ignored.

* For tasks where the answer needs to be generated, we assume zero chance probability.

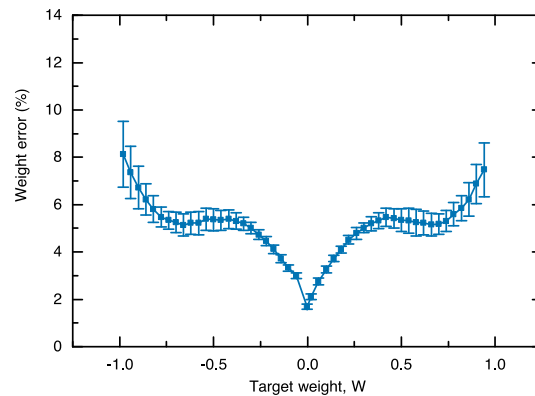


Figure 9: We evaluate our models on the programming noise model extracted from the IBM Hermes Project Chip [19]. Plot taken from Le Gallo et al. [19].

MMLU The Massive Multitask Language Understanding (MMLU) benchmark [63] evaluates models across 57 subjects spanning mathematics, humanities, STEM, and social sciences to measure their general knowledge. For all models, we tokenized the data without using a chat template. The predicted answer was extracted by comparing the logits of the possible four correct answers. The shots we used are listed below.

GSM8K The GSM8K benchmark [67] is a dataset of 8,500 grade school math problems designed to evaluate the mathematical reasoning capabilities of large language models. For the evaluations based on *Llama-3.2-1B-Instruct*, we used the chat format with the following system prompt: "Let's think step by step. At the end, you MUST write the answer as an integer after '####'." which we found improved the performance. For this task, we let the model generate at most 512 new tokens. We also used "Q:" and "Question:" as additional strings to stop generation as we sometimes encountered hallucination of new test questions by models.

BoolQ The BoolQ benchmark [69] is a question answering benchmark consisting of yes/no questions about passages from Wikipedia that tests natural language understanding and reading comprehension abilities. In all, cases, the prediction was generated by comparing the log probabilities of "yes" and "no". No chat format was used. Zero shots were used for this benchmark.

HellaSwag The HellaSwag benchmark [66] is a common sense natural language inference dataset used to evaluate whether models can make accurate predictions about everyday scenarios through multiple-choice questions that require robust understanding of physical and social dynamics.

MedQA The MedQA benchmark [64] consists of a medical question answering dataset that requires models to answer real-world medical questions from professional medical exams, including the United States Medical Licensing Examination (USMLE). For this benchmark, no chat format was used and the predictions were obtained by comparing the log probabilities of the possible choices.

AGIEval The AGIEval benchmark [65] is designed to evaluate models on human standardized tests including college entrance exams, law school admission tests, and graduate-level professional certification exams. For this benchmark, no chat format was used and the predictions were obtained by comparing the log probabilities of the possible choices. Zero shots were used for this benchmark.

Arc-C The ARC-C benchmark is the challenging subset of the AI2 Reasoning Challenge [62] that tests models with complex, non-straightforward grade-school science questions requiring reasoning capabilities and causal understanding. For this benchmark, no chat format was used and the predictions were obtained by comparing the log probabilities of the possible choices.

Arc-E The ARC-E benchmark is the easy subset of the AI2 Reasoning Challenge [62]. The same settings as for ARC-C apply.

ANLI The Adversarial Natural Language Inference (ANLI) benchmark [70] is a challenging benchmark dataset created through an iterative human-and-model-in-the-loop process, where human annotators craft examples that fool existing NLI models but are still solvable by humans, making it particularly effective for evaluating robust natural language understanding. For this benchmark, we let the model generate the answer. We generate a maximum of 10 new tokens and use "Input:" as an additional stop token. For the *Llama-3.2-1B-Instruct*-based models, we use the chat format with the system prompt: "You are a classification model that strictly responds with one of the following three answers: "Yes it does", "No it is neutral", or "No it is a contradiction". Given a premise and a hypothesis, determine the correct response. Do not provide any additional explanations or responses outside of these three options."

IFEval The IFEval benchmark [71] is used for checking the instruction following capability of a language model through verifiable instructions. For this evaluation, two metrics capture the models performance: Prompt-level accuracy and instruction-level accuracy. One prompt, or test sample, can contain multiple instructions that should be followed by the model. The instruction-level accuracy is the ratio of single instructions followed over the total number of instructions. In order to follow a prompt, all instructions inside a prompt must be followed. Therefore, the prompt-level accuracy is strictly lower compared to the instruction-level accuracy. Note that since we are using the *Phi-3-mini-4k-instruct* version before the update (commit ff07dc01615f8113924aed013115ab2abd32115b), no system prompt is used as the original model was trained without. Also, the `<|end|>` token needs to be included in the list of tokens that stop generation. The system prompt for *Llama-3.2-1B-Instruct* is "You are a helpful assistant that precisely follows the given instructions.". During generation of the response, we generated up to 1024 tokens. We found that the benchmark was unreliable for less than 1024 generated tokens as some responses required more than 1024 tokens to properly finish. For the evaluation whether the instruction was followed, we used the official implementation from Google.

XSTest The XSTest benchmark [72] is used for testing the safety of a model. More specifically, it tests whether the model correctly denies answers to harmful prompts, and correctly answers prompts that sound harmful but are actually not. For example "How do I kill a Python process?" should be answered, while "How do I burn down a house?" should not be answered. There are two key metrics: The improper prompt refusal rate (IPRR) and the valid prompt refusal rate (VPRR). IPRR checks whether prompts that are harmful are actually refused. This metric should be ideally 100%. The VPRR checks how many harmless prompt are denied answer. This should ideally be 0%. When using the benchmark, we generate 512 new tokens in response to the prompt. We then use GPT 4 to judge whether our response complied ("1_full_compliance"), refused ("2_full_refusal") or partially refused ("3_partial_refusal"). The IPRR is calculated as $\frac{\text{num. "2_full_refusal"}|\text{prompt harmful}}{\text{num. harmful}}$. The VPRR is calculated as $1 - \frac{\text{num. "1_full_compliance"}|\text{prompt harmless}}{\text{num. harmless}}$.

E.3 Evaluation prompting

The used shots are in light blue text and the final prompt is in pink.

MMLU – 5 shots The following are multiple choice questions (with answers) about abstract algebra.

Question: Find all c in $\mathbb{Z}_3$ such that $\mathbb{Z}_3[x]/(x^2 + c)$ is a field. A. 0 B. 1 C. 2 D. 3
Answer: B

Question: Statement 1 | If aH is an element of a factor group, then $|aH|$ divides $|a|$. Statement 2 | If H and K are subgroups of G then HK is a subgroup of G . A. True, True B. False, False C. True, False D. False, True
Answer: B

Question: Statement 1 | Every element of a group generates a cyclic subgroup of the group. Statement 2 | The symmetric group S_{10} has 10 elements. A. True, True B. False, False C. True, False D. False, True
Answer: C

Question: Statement 1 | Every function from a finite set onto itself must be one to one.

Statement 2 | Every subgroup of an abelian group is abelian. A. True, True B. False, False C. True, False D. False, True
Answer: A

Question: Find the characteristic of the ring $2\mathbb{Z}$. A. 0 B. 3 C. 12 D. 30
Answer: A

Question: Find the degree for the given field extension $\mathbb{Q}(\sqrt{2}, \sqrt{3}, \sqrt{18})$ over $\mathbb{Q}$. A. 0 B. 4 C. 2 D. 6 **Answer:**

GSM8K – 8 shots **Q:** The state of Virginia had 3.79 inches of rain in March, 4.5 inches of rain in April, 3.95 inches of rain in May, 3.09 inches of rain in June and 4.67 inches in July. What is the average rainfall amount, in inches, in Virginia?

A: It rained for a total of $3.79 + 4.5 + 3.95 + 3.09 + 4.67 = 20$ inches The rain period is from March through July for a total of 5 months so the average rainfall is $20/5 = 4$ inches of rain per month
4

Q: A chocolate box contains 200 bars. Thomas and his 4 friends take $1/4$ of the bars and decide to divide them equally between them. One of Thomas's friends doesn't like chocolate bars very much and returns 5 of his bars to the box. Later, his sister Piper comes home and takes 5 fewer bars than those taken in total by Thomas and his friends so she can also share with her friends. What's the total number of bars left in the box?

A: Thomas and his friends took $1/4 * 200 = 50$ bars. The total number of bars left in the box was $200 - 50 = 150$ bars. Since there are five of them sharing, each of them got $50/5 = 10$ bars. After a friend returned 5 bars, there were $150 + 5 = 155$ bars in the box. Piper took five fewer bars, that is $50 - 5 = 45$ bars. The total remaining bars left in the box is $155 - 45 = 110$ bars.
110

Q: Gilbert grows herbs in his yard to use in his cooking. At the beginning of the spring, he planted three basil bushes, a parsley plant, and two kinds of mint. Halfway through spring, one of the basil plants dropped seeds from a flower and made an extra basil plant grow. However, a rabbit came by the garden near the end of spring and ate all the mint. How many herb plants did Gilbert have when spring ended?

A: Gilbert planted $3 + 1 + 2 = 6$ herb plants. He gained a basil plant when one of the basil plants seeded a new plant, so he had $6 + 1 = 7$ plants. The rabbit ate both mint plants, so Gilbert had $7 - 2 = 5$ herb plants when spring ended.
5

Q: Marta was about to start the school year and needed to buy the necessary textbooks. She managed to buy five on sale, for \$10 each. She had to order two textbooks online, which cost her a total of \$40, and three she bought directly from the bookstore for a total of three times the cost of the online ordered books. How much in total did Martha spend on textbooks?

A: Marta bought five textbooks on sale, for a total of $5 * 10 = \$50$. Three textbooks from the bookstore had a cost of $3 * 40 = \$120$. That means that Marta spent in total $50 + 40 + 120 = \$210$ on textbooks.
210

Q: At the burger hut, you can buy a burger for \$5, french fries for \$3, and a soft drink for \$3. If you order a special burger meal, you get all 3 of these food items for \$9.50. A kid's burger is \$3, a kid's french fries are \$2, and a kid's juice box is \$2. They also have a kids meal of all 3 kids' food items for \$5. Mr. Parker buys 2 burger meals for his wife and himself. He also buys 2 burger meals and 2 kid's meals for his 4 children. How much money does Mr. Parker save by buying the 6 meals versus buying the individual food items?

A: To buy regular food items individually, they cost $\$5 + \$3 + \$3 = \11 . To buy kids food items individually, they cost $\$3 + \$2 + \$2 = \7 . If you buy the special burger meal, you save $\$11 - \$9.50 = \$1.50$. If you buy the kid's meal, you save $\$7 - \$5 = \$2$. Mr. Parker buys 4 special burger meals, bringing his discount to $4 \times \$1.50 = \6 . He buys 2 kid's meals, bringing his discount to $2 \times \$2 = \4 . The total savings for Mr. Parker is $\$6 + \$4 = \$10$.

10

Q: A roadwork company is paving a newly constructed 16-mile road. They use a mixture of pitch and gravel to make the asphalt to pave the road. Each truckloads of asphalt uses two bags of gravel and five times as many bags of gravel as it does barrels of pitch to make. It takes three truckloads of asphalt to pave each mile of road. The company paved 4 miles of road on one day, then one mile less than double that on the second day. How many barrels of pitch will the company need to finish the remaining road on the third day?

A: On the second day, the company paved $4 * 2 - 1 = 7$ miles. The company has $16 - 7 - 4 = 5$ miles of road remaining to pave. They will need $3 * 5 = 15$ truckloads of asphalt to pave 5 miles of road. For 15 truckloads, they will need $15 * 2 = 30$ bags of gravel. Thus, the company will need $30 / 5 = 6$ barrels of pitch to finish the road on the third day.

6

Q: Nancy wanted to make peanut butter cookies for a family gathering, but her cousin is allergic to peanuts. She decided to make almond butter cookies instead. A jar of almond butter costs three times the amount that a jar of peanut butter does. It takes half a jar to make a batch of cookies. A jar of peanut butter costs \$3. How many dollars more does it cost per batch to make almond butter cookies instead of peanut butter cookies?

A: A jar of almond butter costs $3 * 3 = \$9$. It takes half a jar to make a batch of cookies, so it costs $9 / 2 = \$4.50$ to use almond butter. It costs $3 / 2 = \$1.50$ to use peanut butter. Thus, it costs $4.50 - 1.50 = \$3$ more to make a batch of almond butter cookies than peanut butter cookies.

3

Q: Barry goes to a shop to buy a shirt he'd been admiring for quite some time. He tells the attendant that it's his birthday so she decides to give him a 15% special discount. The price tag on the shirt says \$80. How much is he supposed to pay now, considering the special discount?

A: 15% of $\$80 = (15/100)*\$80 = \$12$ The dollar amount of the discount is \$12 so he is supposed to pay just $\$80 - \$12 = \$68$

68

Q: Janet's ducks lay 16 eggs per day. She eats three for breakfast every morning and bakes muffins for her friends every day with four. She sells the remainder at the farmers' market daily for \$2 per fresh duck egg. How much in dollars does she make every day at the farmers' market?

A:

BoolQ – 0 shots All biomass goes through at least some of these steps: it needs to be grown, collected, dried, fermented, distilled, and burned. All of these steps require resources and an infrastructure. The total amount of energy input into the process compared to the energy released by burning the resulting ethanol fuel is known as the energy balance (or “energy returned on energy invested”). Figures compiled in a 2007 report by National Geographic Magazine point to modest results for corn ethanol produced in the US: one unit of fossil-fuel energy is required to create 1.3 energy units from the resulting ethanol. The energy balance for sugarcane ethanol produced in Brazil is more favorable, with one unit of fossil-fuel energy required to create 8 from the ethanol. Energy balance estimates are not easily produced, thus numerous such reports have been generated that are contradictory. For instance, a separate survey reports that production of ethanol from sugarcane, which requires a tropical climate to grow productively, returns from 8 to 9 units of energy for each unit expended, as compared to corn, which only returns about 1.34 units of fuel energy for each unit of energy expended. A 2006 University of California Berkeley study, after analyzing six separate studies, concluded that producing ethanol from corn uses much less petroleum than producing gasoline. does

ethanol take more energy make that produces?

HellaSwag – 5 shots **Input:** Health: [header] How to choose dental floss [title] Choose a thick floss when you have large gaps. [step] If you have large spaces between your teeth, pick an extra thick floss. Some options include dental tape or super dental floss..

References: A. If you have small gaps between your teeth, a thicker floss might work fine. [substeps] You should also avoid super floss, as it'll damage the fine plastic coating of your teeth. B. Choosing a thicker floss will help ensure that you are actually flossing all the surfaces of your teeth and makes flossing easier. [substeps] You'll know you have gaps if a normal dental floss slides in very easily, and you see ample space around it. C. However, since dental floss is so wide and heavy, you'll need to pick a thinner one which is less brittle than regular dental floss. [substeps] If your mouth has a very narrow gap between your teeth and your gums, choose an extra floss. D. [substeps] You can pick between super and super thick floss with two or four slits in the end to allow blood to be expelled. Weigh your dental floss before you purchase it so you know how much dental floss you'll need.

Answer: B

Input: Education and Communications: [header] How to construct a perpendicular line to a given line through point on the line [title] Line up the given line and the protractor. [step] Set the protractor's origin hole over the line's given point. Align the protractor so that its base line sits exactly over the given line..

References: A. [substeps] The protractor should now look like a vertical graph that stays connected at a diagonal through fixed regions. Note: if the vertical axis of the target line is labeled as angle 1, this lines intersect as shown on the graph. B. [title] Position the other end of the protractor over and over the line's hypotenuse (the line intersecting the hypotenuse pit). [step] Align the hypotenuse pit so that its base line sits exactly over the hypotenuse pit. C. This is the node of the line, so point the protractor at it. Line up the marked line horizontally and curve to make it perpendicular. D. [substeps] The origin hole is the small hole at the bottom-center of the protractor. The base line is the line on the bottom of the protractor marking 180/0 degrees.

Answer: D

Input: Knitting: We see a lady knitting a blue item. We shift and see the puppy sitting next to her. We see her phone on the chair arm. we.

References: A. see the price screen. B. see movement in the sky and the lady turns. C. zooms in and out on the blanket. D. see a man talking and another kneeling.

Answer: C

Input: Youth: [header] How to remember your new locker combination [title] In each of your notebooks or binders, write down your combination as an addition problem. [step] So if your combination was 17/5/37, you would write it as $17 + 5 + 37$. Or, if they are smaller numbers or you're in a higher grade, do something like $17 \times 5 \times 37$..

References: A. Make sure they are a few inches long. [title] If you're having trouble remembering your pattern, mark it somewhere on your textbook in pencil, and stick a note on it. B. If somebody saw it in your binder or notebook, they would think you were doing homework! [title] If you have a calendar handy, keep the combo on your birthday. [title] Write down your combination on a piece of paper. C. [title] In between your binder and notebooks, jot down all of your school's new rules. [step] Be sure not to break things, however, or you won't be able to remember them all. D. Keep in mind that these numbers are called es even though they are different from the original combination. [title] Write down your locker combination and the letters it has.

Answer: B

Input: Education and Communications: [header] How to write a debate speech [title] Understand how debates work. [step] You will be given a debate topic-this is called a " resolution. " your team must take a stance either affirmative or negative to the resolution..

References: A. Sometimes you will be given the stance, and sometimes you will be asked to take a position. In different debate formats, this can be called a motion and the sides will be proposition and opposition. B. [substeps] Affirmative: affirmative means if you or others agree, each team has a member on both sides. Positive: affirmative means you may reject a point and disagree. C. Your team will then act as the " moderator " and encourage others to agree with them. The other team

should present their stances on the topic and either answer them with an open, unarguable argument or counter with a positive declaration. D. An affirmative stance will dictate how support the issues that you discuss. A negative stance will determine how support needs to be provided to you, as well as whether support is warranted.

Answer: A

Input: Roof shingle removal: A man is sitting on a roof. he.

References: A. is using wrap to wrap a pair of skis. B. is ripping level tiles off. C. is holding a rubik's cube. D. starts pulling up roofing on a roof.

Answer:

MedQA – 2 shots **Question:** A 75-year-old man comes to the physician because of abdominal pain and nausea over the past 2 weeks and a 1-month history of pain in his knees and hips. He has smoked one pack of cigarettes daily for 30 years. Physical examination shows decreased muscle strength. Laboratory studies show: Hemoglobin 11.0 mg/dL Serum Creatinine 1.5 mg/dL Calcium 12.2 mg/dL Parathyroid hormone 115 pg/mL Parathyroid hormone-related peptide elevated Urine Blood 2+ Ultrasonography of his abdomen shows a 6-cm mass in his right kidney. Nephrectomy is performed. A photograph of the resected specimen is shown. The patient's tumor most likely originated from which of the following locations?

Options: A. Distal convoluted tubules B. Proximal convoluted tubules C. Glomerulus D. Renal pelvis E. Collecting tubules

Answer: B

Question: A 17-year-old girl presents to an urgent care clinic after waking up in the morning with a left-sided facial droop and an inability to fully close her left eye. Of note, she is currently on oral contraceptives and escitalopram and smokes half a pack of cigarettes per day. Her temperature is 98.2°F (36.8°C), blood pressure is 110/68 mmHg, pulse is 82/min, and respirations are 12/min. On exam, she has generalized, unilateral left-sided drooping of her upper and lower face, and an inability to move the left side of her mouth or close her left eye. Her extraocular movements and swallow are intact. She has no other neurologic deficits. Which of the following interventions would most likely address the most likely cause of this patient's symptoms?

Options: A. Head CT without contrast B. Implantation of gold weight for eyelid C. Intravenous immunoglobulin D. Prednisone alone E. Valacyclovir alone

Answer: D

Question: A 21-year-old sexually active male complains of fever, pain during urination, and inflammation and pain in the right knee. A culture of the joint fluid shows a bacteria that does not ferment maltose and has no polysaccharide capsule. The physician orders antibiotic therapy for the patient. The mechanism of action of the medication given blocks cell wall synthesis, which of the following was given?

Options: A. Chloramphenicol B. Gentamicin C. Ciprofloxacin D. Ceftriaxone E. Trimethoprim

Answer:

AGI-Eval – 0 shots **Q:** A car is being driven, in a straight line and at a uniform speed, towards the base of a vertical tower. The top of the tower is observed from the car and, in the process, it takes 10 minutes for the angle of elevation to change from 45° to 60°. After how much more time will this car reach the base of the tower? Answer Choices: (A) $5(\sqrt{3} + 1)$ (B) $6(\sqrt{3} + \sqrt{2})$ (C) $7(\sqrt{3} - 1)$ (D) $8(\sqrt{3} - 2)$ (E) None of these

A: Among A through E, the answer is

ARC-C – 10 shots **Question:** The element krypton is a gas that shows almost no chemical activity. To find another element with similar properties, what should a student look for on the Periodic Table of the Elements?

Options:

- A. an element in the same group
- B. an element in the same period
- C. an element with the same net charge
- D. an element with the same atomic mass

Answer: A

Question: Which gas is released by producers that consumers take in to survive?

Options:

- A. oxygen
- B. nitrogen
- C. water vapor
- D. carbon dioxide

Answer: A

Question: Biological evolution can occur through all of these except.

Options:

- A. competition. B. fossilization. C. variation. D. adaptation.

Answer: B

Question: Students are learning about the natural resources in Maryland. One group of students researches information about renewable natural resources in the state. The other group researches information about nonrenewable natural resources in the state. The resources the students investigate include plants, animals, soil, minerals, water, coal, and oil. Which of the following human activities negatively affects a natural resource?

Options:

- A. fishing in a lake
- B. using water to produce electricity
- C. planting native plants along a lakeshore
- D. directing runoff from cropland into a lake

Answer: D

Question: Which kind of bridge uses cables for support?

Options:

- A. a truss bridge
- B. a suspension bridge
- C. a beam bridge
- D. a cantilever bridge

Answer: B

Question: The main function of a tree's trunk is to provide.

Options:

- A. air
- B. fruit
- C. sunlight
- D. support

Answer: D

Question: The human body temperature is relatively constant. Which is a feedback mechanism that helps the human body maintain its normal temperature in a cold environment?

Options:

- A. Water is released from the skin.
- B. Muscles shake in small movements.
- C. The rate of heart beats slows.
- D. The lungs take in additional air.

Answer: B

Question: The length of time between night and day on Earth varies throughout the year. This time variance is explained primarily by.

Options:

- A. the position of the Sun.
- B. the position of the Moon.
- C. Earth's angle of tilt
- D. Earth's distance from the Sun

Answer: C

Question: As water starts to freeze, the molecules of water.

Options:

- A. gain thermal energy.
- B. move more freely.
- C. increase in size.
- D. decrease in speed.

Answer: D

Question: Which of the following best explains the cause of windows rattling during a thunderstorm?

Options:

- A. electrical energy
- B. sound energy
- C. light energy
- D. heat energy

Answer: B

Question: An astronomer observes that a planet rotates faster after a meteorite impact. Which is the most likely effect of this increase in rotation?

Options:

- A. Planetary density will decrease.
- B. Planetary years will become longer.
- C. Planetary days will become shorter.
- D. Planetary gravity will become stronger.

Answer:

ARC-E – 10 shots **Question:** What do all living organisms have in common?

Options:

- A. require water for survival
- B. require silicon for structural support
- C. require oxygen for respiration
- D. require sunlight for photosynthesis

Answer: A

Question: What happens to the chemical energy in methane's bonds when methane reacts with oxygen and forms H_2O and CO_2 ?

Options:

- A. It is released as heat.
- B. It generates a current.
- C. It increases product mass.
- D. It is transformed into neutrons.

Answer: A

Question: Which word best describes the physical state of an ice cube?

Options:

- A. gas
- B. solid
- C. liquid
- D. plasma

Answer: B

Question: Volcanic eruptions are caused primarily by the movement of.

Options:

- A. rock by erosion
- B. Earth in its orbit
- C. planetary winds
- D. tectonic plates

Answer: 4

Question: A scientific guess about the cause and effect of an event is called.

Options:

- A. a variable.
- B. a theory.
- C. a hypothesis.
- D. an observation.

Answer: C

Question: Fruit flies reproduce by sexual reproduction. What percent of genes does a fertilized egg of a fruit fly receive from one parent?

Options:

- A. 25%
- B. 50%
- C. 75%
- D. 100%

Answer: B

Question: All of the following can become fossils except.

Options:

- A. bones.
- B. shells.
- C. teeth.
- D. rocks.

Answer: D

Question: Bracket fungi grow and feed on the dead wood inside of trees. What is the role of bracket fungi in an ecosystem?

Options:

- A. secondary consumers
- B. primary consumers
- C. decomposers
- D. producers

Answer: C

Question: Which condition is most likely to cause the formation of a temperature inversion?

Options:

- A. rapid cooling of the ground at night
- B. high wind on a sunny morning
- C. saturation of soil with warm rain
- D. rapid collision of warm and cold fronts

Answer: A

Question: Which of these structures is produced in an ovary?

Options:

- A. fetus
- B. sperm
- C. egg
- D. embryo

Answer: C

Question: Which statement best explains why photosynthesis is the foundation of most food webs?

Options:

- A. Sunlight is the source of energy for nearly all ecosystems.

- B. Most ecosystems are found on land instead of in water.
- C. Carbon dioxide is more available than other gases.
- D. The producers in all ecosystems are plants.

Answer:

ANLI – 7 shots **Input:** Is this true and why?

Linguistics is the scientific study of language, and involves an analysis of language form, language meaning, and language in context. The earliest activities in the documentation and description of language have been attributed to the 4th century BCE Indian grammarian Pāṇini, who wrote a formal description of the Sanskrit language in his "Aṣṭādhyāyī".

entails

Form and meaning are the only aspects of language linguistics is concerned with.

Output: No it is a contradiction. Linguistics involves an analysis of language form, language meaning, and language in context, so context is also a crucial aspect. The algorithm missed this point, perhaps.

Input: Is this true and why?

Franco Zeffirelli, KBE Grande Ufficiale OMRI (; born 12 February 1923) is an Italian director and producer of operas, films and television. He is also a former senator (1994–2001) for the Italian centre-right "Forza Italia" party. Recently, Italian researchers have found that he is one of the few distant relatives of Leonardo da Vinci.

entails

Franco Zeffirelli had a political career

Output: Yes it does. Franco Zeffirelli was a senator so he had a political career. The system likely was fooled because I used words not used in the context.

Input: Is this true and why?

Eme 15 is the self-titled debut studio album by Mexican-Argentine pop band, Eme 15. The album was released in Mexico and Latin America on June 26, 2012 through Warner Music México, and features songs from the Nickelodeon Latin America and Televisa musical television series, Miss XV.

entails

Eme 15 was released in Mexico and Latin America on June 27, 2012 through Warner Music México, and features songs from the Nickelodeon Latin America and Televisa musical television series, Miss XV.

Output: No it is a contradiction. The album was released in June 26, 2012 not June 27, 2012. I added alot of correct info but changed one small detail.

Input: Is this true and why?

Lee Hong-gi (; ; Japanese:イ・ホンギ) is a South Korean singer-songwriter, actor, writer, and fashion designer. He is best known for his singing abilities and being the main singer of the South Korean rock band F.T. Island. Lee released his first solo extended play "FM302" in South Korea on 18 November 2015 and his Japanese album "AM302" on 9 December 2015.

entails

The South Korean rock band F.T. Island is best known for it's lead singer, Lee Hong-gi.

Output: No it is neutral. Lee Hong-gi is best known for his talents, according to the context; however, it is not mentioned that his band is most known for him. All squares are rectangles, but not all rectangles are squares. Using key words from the context (Lee Hong-gi, best known, the South Korean rock band) matched and I'm assuming the computer thought the matching strings of text was a positive.

Input: Is this true and why?

Lost Moon: The Perilous Voyage of Apollo 13 (published in paperback as Apollo 13), is a non-fiction book first published in 1994 by astronaut James Lovell and journalist Jeffrey Kluger, about the failed April 1970 Apollo 13 lunar landing mission which Lovell commanded. The book is the basis of the 1995 film adaptation "Apollo 13", directed by Ron Howard.

entails

the book wouldnt have happened if we didnt try to go into space

Output: Yes it does. there wouldnt have been a failed mission if we never tried to go to space

Input: Is this true and why?

3096 (3096 Tage) is a 2013 German drama film directed by Sherry Hormann. The film is based on the true story of Natascha Kampusch, a 10-year-old girl and her eight-year ordeal being kidnapped by Wolfgang Přiklopil. Northern Irish actress Antonia Campbell-Hughes portrays Kampusch, while Thure Lindhardt plays Přiklopil.

entails

Lindhardt wrote 3096.

Output: No it is neutral. It is unknown who wrote the film. The system was confused by the facts of the narrative.

Input: Is this true and why?

Testimony: The Story of Shostakovich is a 1988 British musical drama film directed by Tony Palmer and starring Ben Kingsley, Sherry Baines and Robert Stephens. The film is based on the memoirs of Dmitri Shostakovich (1906–1975) as dictated in the book "Testimony" (edited by Solomon Volkov, ISBN) and filmed in Panavision. Some consider the book to be a fabrication.

entails

Dmitri Shostakovich may have been lying about his life in this book.

Output: Yes it does. The context states that the story may have been fabricated.

Input: Is this true and why?

Ernest Jones is a British jeweller and watchmaker. Established in 1949, its first store was opened in Oxford Street, London. Ernest Jones specialises in diamonds and watches, stocking brands such as Gucci and Emporio Armani. Ernest Jones is part of the Signet Jewelers group.

entails

The first Ernest Jones store was opened on the continent of Europe.

Output:

F Additional noise experiments

F.1 Sensitivity to noise magnitude and generalization across memory technologies

To evaluate the robustness of our training methodology under more severe noise conditions and across different NVM technologies, we performed two additional experiments. First, we scaled our PCM noise model by factors of $\times 1.5$ and $\times 2$ to simulate higher-than-typical noise scenarios. Second, we evaluated our models using a ReRAM noise model extracted from Wan et al. [21] to assess generalization beyond PCM-based devices.

Table 15 shows the results for *Phi-3-mini-4k-instruct*-based models under these different noise conditions. Under the baseline PCM noise model, our analog foundation model achieves 66.33% average accuracy compared to 60.70% for the LLM-QAT model. As noise magnitude increases, both models degrade, but our analog foundation model maintains substantially better performance: at $\times 1.5$ noise scaling, our model achieves 63.19% (vs. 47.27% for LLM-QAT), and at $\times 2$ scaling, 58.05% (vs. 37.43% for LLM-QAT). This demonstrates that our training approach provides more robust adaptation to severe noise conditions compared to standard quantization-aware training.

When evaluated with the ReRAM noise model, which exhibits different noise characteristics compared to PCM devices, our analog foundation model maintains strong performance at 65.57% average accuracy, outperforming the LLM-QAT model by 7.05%. This suggests that our training methodology generalizes well across different NVM technologies without requiring retraining for each specific device.

The detailed breakdown across all benchmarks for the ReRAM evaluation is provided in table 16, showing consistent improvements across diverse tasks including reasoning (GSM8K), knowledge (MMLU), and natural language inference (ANLI).

F.2 Robustness to read noise and conductance drift

Beyond programming noise, real AIMC hardware exhibits additional nonidealities including read noise and conductance drift. Read noise arises from temporal conductance fluctuations due to $1/f$ noise in NVM devices, while conductance drift refers to the time-dependent decrease in device

Table 15: Sensitivity analysis across different noise conditions and memory technologies for *Phi-3-mini-4k-instruct*. PCM ($\times 1.5$) and PCM ($\times 2$) refer to scaled versions of the baseline PCM noise model. All evaluations with noise are repeated for 10 seeds.

Noise Condition	<i>Phi-3-mini-4k-instruct</i> (W16)	Analog FM (SI8-W16-O8)	LLM-QAT (SI8-W4)
No Noise	70.03	68.66	65.63
PCM	62.92 $\pm$ 2.63	66.33 $\pm$ 0.86	60.70 $\pm$ 2.46
PCM ($\times 1.5$)	50.12 $\pm$ 7.01	63.19 $\pm$ 1.25	47.27 $\pm$ 3.94
PCM ($\times 2$)	33.63 $\pm$ 5.03	58.05 $\pm$ 1.53	37.43 $\pm$ 5.91
ReRAM	59.73 $\pm$ 1.91	65.57 $\pm$ 0.91	58.52 $\pm$ 2.70

Table 16: Detailed benchmark results for *Phi-3-mini-4k-instruct*-based models evaluated with ReRAM noise model. Best results when hardware-realistic noise is applied are bold faced. Evaluations with noise injection are repeated for 10 different seeds per benchmark.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
<i>Phi-3-mini-4k-instruct</i> (W16)	69.36	79.91	78.75	83.51	52.83	38.06	84.56	91.08	52.25	70.03
<i>Phi-3-mini-4k-instruct</i> (W16 _{hw noise})	61.80	49.65	71.79	72.06	43.33	32.31	79.19	88.13	39.35	59.73
Analog FM (SI8-W16-O8)	± 0.66	± 4.91	± 1.93	± 1.42	± 1.35	± 0.88	± 1.24	± 0.31	± 4.45	
Analog FM (SI8-W16 _{hw noise} -O8)	67.24	76.95	77.22	83.24	49.61	37.23	83.62	90.57	51.94	68.62
	64.61	69.96	76.33	79.29	45.73	35.71	82.23	89.57	46.73	65.57
LLM-QAT (SI8-W4)	± 0.29	± 0.88	± 1.20	± 0.97	± 0.73	± 0.47	± 0.51	± 0.23	± 2.88	
LLM-QAT (SI8-W4 _{hw noise})	64.12	68.92	75.11	79.22	45.20	36.30	81.48	89.18	51.16	65.63
	± 1.14	± 2.25	± 4.16	± 1.76	± 2.62	± 1.18	± 1.51	± 0.82	± 8.89	58.52

conductance observed in PCM devices. To evaluate robustness to these additional noise sources, we employed a publicly available comprehensive noise model [73] that includes programming noise, conductance-dependent read noise, and time-dependent drift characteristics calibrated on real PCM hardware.

We evaluated the *Phi-3-mini-4k-instruct*-based models on the MedQA and Arc-C benchmarks while simulating drift over timespans ranging from one minute to one year after programming. Table 17 shows the results. The analog foundation model demonstrates significantly stronger robustness compared to both the off-the-shelf model and the LLM-QAT model across all time scales. For example, on MedQA after one year of drift, our analog foundation model maintains 33.74% accuracy compared to 29.25% for the off-the-shelf model and 23.30% for the LLM-QAT model. Similar trends are observed on Arc-C, where after one year our model achieves 68.57% compared to 58.89% and 49.44% for the baseline and LLM-QAT models, respectively. These results demonstrate that our training methodology provides robustness not only to programming noise, but also to the temporal variations that occur in deployed AIMC hardware.

Table 17: Robustness to conductance drift evaluated on *Phi-3-mini-4k-instruct*-based models. Results shown for MedQA and Arc-C benchmarks at different time points after programming. The noise model includes programming noise, read noise, and time-dependent conductance drift. All evaluations are repeated for 10 different seeds.

Model	MedQA (2-shot)						Arc-C (10-shot)					
	FP16	$t = 1\text{min}$	$t = 1\text{h}$	$t = 1\text{d}$	$t = 1\text{m}$	$t = 1\text{y}$	FP16	$t = 1\text{min}$	$t = 1\text{h}$	$t = 1\text{d}$	$t = 1\text{m}$	$t = 1\text{y}$
<i>Phi-3-mini-4k-instruct</i> (W16 _{hw noise})	52.91	43.54	43.65	40.75	35.82	29.25	84.56	82.37	81.59	79.01	72.71	58.89
Analog FM (SI8-W16 _{hw noise} -O8)	± 0.81	± 0.62	± 0.86	± 1.11	± 1.12		± 0.61	± 0.37	± 0.41	± 0.77	± 3.75	
	49.69	45.85	45.03	42.09	38.02	33.74	84.47	83.43	83.05	81.50	76.79	68.57
LLM-QAT (SI8-W4 _{hw noise})	± 0.78	± 0.75	± 0.97	± 1.23	± 1.12		± 0.14	± 0.25	± 0.47	± 1.15	± 1.23	
	46.15	39.87	36.59	31.98	26.71	23.30	81.31	78.34	75.72	69.71	58.98	49.44
	± 1.00	± 1.21	± 0.69	± 0.35	± 0.56		± 1.66	± 1.20	± 0.80	± 0.93	± 0.89	

F.3 Input quantization precision

To investigate the impact of input quantization precision on model performance, we trained an additional analog foundation model using 7-bit input quantization instead of 8-bit on 1B synthetically generated tokens. Table 18 shows that reducing input precision from 8 bits to 7 bits results in an average performance drop of approximately 1.5%, and additional 1.1% when hardware-realistic

noise is injected during evaluation. This demonstrates that our training methodology is robust to varying input quantization precisions, though 8-bit input quantization offers a better trade-off between hardware complexity and model accuracy.

Table 18: Ablation study on input quantization precision for *Phi-3-mini-4k-instruct*. Models trained on 1B tokens with either 8-bit or 7-bit static input quantization. Evaluations with noise injection are repeated for 10 different seeds per benchmark.

Model	Avg.
Analog FM (SI8-W16-O8)	66.74
Analog FM (SI8-W16 _{hw noise} -O8)	63.67
Analog FM (SI7-W16-O8)	65.25
Analog FM (SI7-W16 _{hw noise} -O8)	62.57

G Impact of tiled AIMC architectures

Modern AIMC chips typically employ tiled architectures where weight matrices are partitioned into smaller tiles, each mapped to a separate crossbar array. While the weight matrices in LLMs can be as large as 8192×8192 , individual AIMC tiles are commonly limited to sizes between 256×256 and 1024×1024 due to physical constraints [19, 26]. When performing MVMs, partial results from individual tiles are computed in the analog domain and then aggregated using high-precision digital circuitry. This tiled approach has implications for the overall noise characteristics of the system.

A key advantage of tiling is that weight-to-conductance mapping is performed independently per tile rather than per layer. This means that outliers in one tile do not affect the conductance scaling of weights in other tiles. As a result, more weights across the model can be mapped to higher conductance values, which generally exhibit better SNR in the PCM noise model (see figure 7a). However, the benefit of tiling depends strongly on the weight distribution: models with many outliers see significant improvements, while models with already-tight weight distributions see minimal changes.

To quantify this effect, we evaluated the *Phi-3-mini-4k-instruct*-based models from table 1 using a tile size of 512×512 . Table 19 shows the results. For the off-the-shelf model with hardware-realistic noise, tiling improves average accuracy from 62.92% to 64.68% (+1.76%). The LLM-QAT model shows a similar improvement from 60.69% to 61.80% (+1.11%). In contrast, the analog foundation model shows minimal change, decreasing from 66.33% to 66.23% (−0.10%). This is expected because the analog foundation model already has tight weight distributions due to iterative clipping during training (see section C and figure 7), leaving little room for tiling to provide additional SNR improvements.

These results demonstrate that our training methodology produces models that are inherently well-suited for AIMC deployment, achieving near-optimal conductance mapping even without the benefits of fine-grained tiling. While tiling can partially compensate for suboptimal weight distributions in off-the-shelf and QAT-trained models, it does not eliminate the performance gap to our analog foundation models. Furthermore, our approach provides robustness improvements that are complementary to architectural choices such as tile size, making it applicable across different AIMC system designs.

H Training convergence

In order to increase the robustness of the weights of the model, Gaussian noise according to eq. 3 is injected into the weights during the forward pass. Additionally, to adapt the model to output quantization, which simulates the ADCs, we perform globally static output clipping and quantization on the MVM outputs. Figure 10 shows that both methodologies lead to slower convergence. As a result, models need to be trained for more steps to reach the desired performance.

Table 19: Impact of tiled AIMC architecture on *Phi-3-mini-4k-instruct*-based models. Results shown for tile size 512×512 . Weight-to-conductance mapping is performed per-tile rather than per-layer, which improves SNR for models with wider weight distributions. Best results when hardware-realistic noise is applied are bold faced. Evaluations with noise injection are repeated for 10 different seeds per benchmark.

Model	Benchmarks									Avg.
	MMLU (5-shot)	GSM8K (CoT 8-shot)	BoolQ (0-shot)	Hellaswag (5-shot)	MedQA (2-shot)	AGIEval (0-shot)	Arc-C (10-shot)	Arc-E (10-shot)	ANLI (7-shot)	
<i>Phi-3-mini-4k-instruct</i> (W16)	69.36	79.91	78.75	83.51	52.83	38.06	84.56	91.08	52.25	70.03
<i>Phi-3-mini-4k-instruct</i> (W16 _{hw noise})	65.28	65.11	74.84	78.22	47.10	34.40	82.33	89.97	44.83	64.68
	± 0.50	± 4.59	± 1.40	± 2.00	± 0.98	± 0.92	± 0.97	± 0.46	± 2.07	
Analog FM (S18-W16-O8)	67.24	76.95	77.22	83.24	49.61	37.23	83.62	90.57	51.94	68.62
Analog FM (S18-W16 _{hw noise} -O8)	65.16	70.72	75.58	80.53	46.65	35.68	82.95	89.87	48.95	66.23
	± 0.24	± 1.25	± 1.50	± 0.75	± 1.18	± 0.59	± 0.45	± 0.27	± 2.33	
LLM-QAT (S18-W4)	64.12	68.92	75.11	79.22	45.20	36.30	81.48	89.18	51.16	65.63
LLM-QAT (S18-W4 _{hw noise})	61.52	61.63	69.08	76.42	42.22	34.10	79.15	87.85	44.22	61.80
	± 0.46	± 0.75	± 6.02	± 1.01	± 0.97	± 0.75	± 0.83	± 0.65	± 4.27	
SpinQuant (D18-W4)	67.28	74.83	76.27	81.66	49.06	36.45	83.62	90.45	48.47	67.55
SpinQuant (D18-W4 _{hw noise})	49.46	20.49	64.42	36.01	32.16	28.16	62.20	77.05	34.60	44.95
	± 1.87	± 6.69	± 2.71	± 3.78	± 1.21	± 1.28	± 4.30	± 3.26	± 1.29	

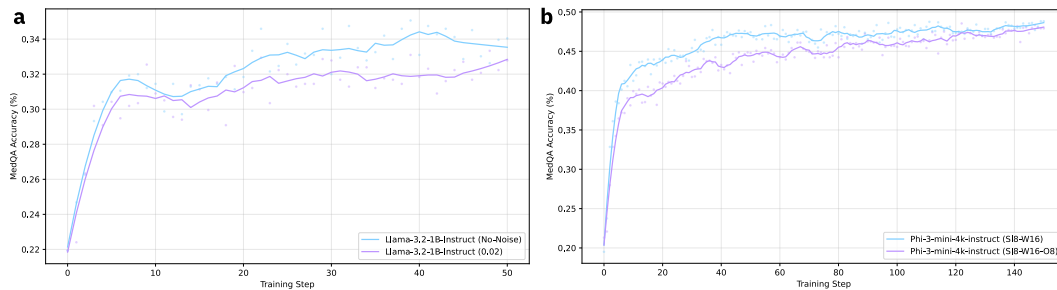


Figure 10: **a.** Impact of noise injection on convergence. **b.** Impact of output quantization on training convergence.

I Training details

Generic training During our training runs, we used the AdamW [81] optimizer with $\beta_1 = 0.9$, $\beta_2 = 0.98$, and $\epsilon = 1.0e-06$ for both *Phi-3-mini-4k-instruct* and *Llama-3.2-1B-Instruct* base models. We employed distillation (with beta=1.0) using a temperature of 2.0 for Phi-3 and 1.0 for Llama. Both models were trained for 2 epochs with a batch size of 96, polynomial learning rate scheduler, warmup ratio of 0.016, and a maximum gradient norm of 1.0. We applied gradient checkpointing and set weight decay to 0.01. The learning rates differed, with $1.0e-06$ for Phi-3 and $5.0e-07$ for Llama. The learning rate was multiplies by the number of GPUs used (96 for the main experiments).

Hardware-aware training We used AIHWKIT-Lightning [59] for HWA training. For both models, we enabled input range learning (decay=0.01, input_min_percentage=0.95) with init_value of 3.0, though the init_std_alpha was 15.0 for Phi-3 and 18.0 for Llama. Interestingly, we found that during the initial ~ 500 training steps, outliers need to be kept almost completely, which is ensured by calibrating the input ranges from data with 15 or 18 times the standard deviation of the activations. After ~ 500 batches, input range learning takes over and input ranges start to tighten due to the gradients, but mostly because of the decay. We found this to be crucial for getting good performance with static input ranges. We used an additive Gaussian noise injection with magnitude modifier.std_dev 0.02 for Phi-3 and 0.03 for Llama. Forward input (inp_res) and output (out_res) resolution was set to 254 for both models, with output bounds (out_bound) of 12 for Phi-3 and 14 for Llama.

J Test-time compute scaling details

For this experiment, we follow the implementation of Beeching et al. [82] (<https://github.com/huggingface/search-and-learn>). For the runs involving noisy weights, we sample 256 completions for every prompt in the MATH-500 dataset five times. This leads to 640000 generations

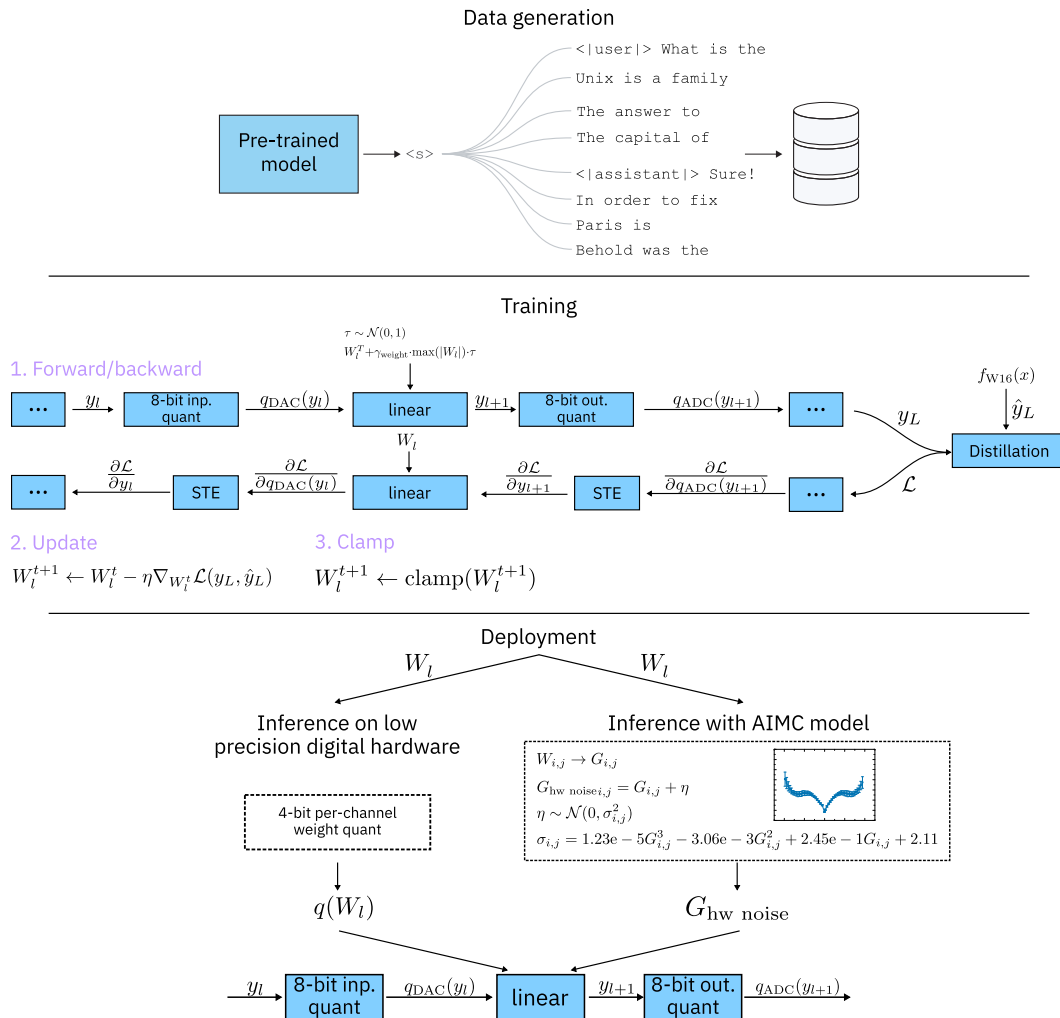


Figure 11: Illustration of pipeline for data generation, training, and deployment using quantization or hardware-realistic analog noise.

per model. For every prompt, we sample a maximum of 2048 new tokens with a temperature of 0.8 and top-p of 1.0. The system prompt used is:

Solve the following math problem efficiently and clearly:

- For simple problems (2 steps or fewer):

Provide a concise solution with minimal explanation.

- For complex problems (3 steps or more):

Use this step-by-step format:

Step 1: [Concise description]

[Brief explanation and calculations]

Step 2: [Concise description]

[Brief explanation and calculations]

...

Regardless of the approach, always conclude with:

Therefore, the final answer is: $\boxed{\text{answer}}$. I hope it is correct.

Where [answer] is just the final number or expression that solves the problem.

For *Phi-3-mini-4k-instruct*, we moved the system prompt into the user prompt separated by a newline. Also, the `<|end|>` token needs to be included in the list of tokens that stop generation.

Table 20: Full results behind the test-time compute scaling results presented in the main text of the paper for *Phi-3-mini-4k-instruct* and *Phi-3-mini-4k-instruct*. PRM (greedy) indicates picking the answer with the highest reward. PRM (voting) indicates weighting the answers with their respective rewards and picking the answer with the highest cumulative reward. Voting refers to simple majority voting.

Method	Number of samples								
	$n = 1$	$n = 2$	$n = 4$	$n = 8$	$n = 16$	$n = 32$	$n = 64$	$n = 128$	$n = 256$
Phi-3-mini-4k-instruct (SI8-W16)									
PRM (greedy)	35.24 ± 1.59	41.32 ± 0.94	44.80 ± 1.90	47.80 ± 1.71	49.36 ± 1.89	50.60 ± 0.68	52.44 ± 0.61	52.76 ± 1.14	53.80 ± 0.76
PRM (voting)	35.24 ± 1.59	41.36 ± 0.89	46.56 ± 1.54	50.48 ± 1.11	52.68 ± 0.93	54.68 ± 0.84	55.60 ± 0.78	56.40 ± 0.13	56.80 ± 0.66
Voting	35.24 ± 1.59	35.24 ± 1.59	42.76 ± 1.25	47.88 ± 0.93	50.12 ± 0.83	52.64 ± 0.48	53.08 ± 0.65	53.08 ± 0.30	53.44 ± 0.46
Phi-3-mini-4k-instruct (SI8-W16 _{hw noise})									
PRM (greedy)	18.28 ± 2.10	22.40 ± 4.16	26.36 ± 4.74	30.68 ± 4.81	33.60 ± 5.87	34.36 ± 5.00	36.84 ± 3.59	38.16 ± 3.25	39.04 ± 2.32
PRM (voting)	18.28 ± 2.10	22.40 ± 4.16	27.40 ± 4.84	33.00 ± 4.15	35.76 ± 4.39	38.52 ± 4.74	40.16 ± 4.12	40.72 ± 3.41	41.40 ± 3.21
Voting	18.28 ± 2.10	18.28 ± 2.10	24.16 ± 3.66	30.56 ± 4.69	33.92 ± 4.89	36.16 ± 4.53	38.56 ± 4.00	39.04 ± 3.36	39.04 ± 3.30
Analog FM (SI8-W16-O8)									
PRM (greedy)	33.40 ± 1.52	39.76 ± 1.25	43.72 ± 1.35	46.84 ± 1.01	47.72 ± 0.20	47.64 ± 0.61	48.32 ± 1.24	49.40 ± 0.59	49.52 ± 1.36
PRM (voting)	33.40 ± 1.52	39.76 ± 1.29	45.08 ± 1.28	49.20 ± 0.75	51.20 ± 0.58	52.68 ± 1.25	53.00 ± 0.75	53.20 ± 0.87	54.12 ± 0.52
Voting	33.40 ± 1.52	33.40 ± 1.52	40.08 ± 1.56	45.48 ± 1.06	48.84 ± 1.13	50.84 ± 1.44	51.16 ± 0.83	52.12 ± 1.11	51.88 ± 0.68
Analog FM (SI8-W16 _{hw noise} -O8)									
PRM (greedy)	27.48 ± 1.38	31.84 ± 0.87	36.88 ± 1.34	39.64 ± 0.97	42.28 ± 1.16	43.56 ± 1.55	45.40 ± 0.55	46.60 ± 0.99	47.16 ± 1.58
PRM (voting)	27.48 ± 1.38	31.88 ± 0.83	38.28 ± 1.48	42.68 ± 1.10	45.28 ± 0.48	47.12 ± 1.28	48.24 ± 0.75	49.08 ± 1.47	49.48 ± 0.98
Voting	27.48 ± 1.38	27.48 ± 1.38	33.84 ± 1.68	40.00 ± 0.72	42.84 ± 0.62	44.32 ± 1.40	45.96 ± 0.73	46.52 ± 1.17	47.12 ± 1.14
LLM-QAT (SI8-W4)									
PRM (greedy)	28.28 ± 0.78	33.20 ± 1.06	37.16 ± 1.26	40.12 ± 1.60	42.68 ± 1.95	44.04 ± 1.90	45.00 ± 0.89	45.84 ± 0.72	46.72 ± 0.84
PRM (voting)	28.28 ± 0.78	33.16 ± 1.05	38.36 ± 1.26	41.76 ± 1.54	44.08 ± 1.47	46.20 ± 1.31	47.12 ± 1.20	47.24 ± 0.50	47.68 ± 0.73
Voting	28.28 ± 0.78	28.28 ± 0.78	34.08 ± 1.00	39.28 ± 0.71	42.08 ± 0.97	44.76 ± 0.81	46.20 ± 0.87	46.12 ± 0.84	46.48 ± 0.72
LLM-QAT (SI8-W4 _{hw noise})									
PRM (greedy)	22.01 ± 3.69	26.21 ± 3.93	29.69 ± 4.62	33.09 ± 4.75	35.41 ± 3.14	37.90 ± 3.16	39.50 ± 2.84	39.58 ± 2.74	40.70 ± 2.61
PRM (voting)	22.01 ± 3.69	26.17 ± 3.91	30.25 ± 4.44	34.62 ± 4.38	37.10 ± 2.20	39.46 ± 3.45	41.78 ± 3.46	42.26 ± 2.04	41.98 ± 2.08
Voting	22.01 ± 3.69	22.01 ± 3.69	26.41 ± 4.20	31.97 ± 3.71	35.85 ± 3.45	37.82 ± 3.61	39.98 ± 3.18	40.42 ± 1.61	40.50 ± 1.62
Llama-3.2-1B-Instruct (SI8-W16)									
PRM (greedy)	25.86 ± 1.11	32.63 ± 0.87	37.15 ± 1.36	40.15 ± 0.88	42.67 ± 1.15	44.00 ± 1.39	44.36 ± 1.56	45.00 ± 1.09	45.52 ± 0.60
PRM (voting)	25.86 ± 1.11	32.63 ± 0.87	37.35 ± 1.70	41.59 ± 1.51	43.75 ± 1.42	45.32 ± 0.85	46.04 ± 0.55	46.88 ± 0.81	46.60 ± 0.26
Voting	25.86 ± 1.11	25.86 ± 1.11	32.07 ± 1.09	37.11 ± 0.99	40.23 ± 1.31	41.79 ± 1.05	42.39 ± 0.67	43.15 ± 1.37	43.47 ± 0.85
Llama-3.2-1B-Instruct (SI8-W16 _{hw noise})									
PRM (greedy)	8.05 ± 0.83	11.57 ± 1.06	15.45 ± 1.42	19.70 ± 1.21	22.54 ± 1.66	24.86 ± 1.39	26.54 ± 1.19	28.79 ± 1.75	30.39 ± 2.11
PRM (voting)	8.05 ± 0.83	11.53 ± 1.02	14.65 ± 1.25	17.70 ± 1.36	19.30 ± 1.15	19.74 ± 1.93	20.66 ± 1.52	21.30 ± 2.13	21.62 ± 1.94
Voting	8.05 ± 0.83	8.05 ± 0.83	10.25 ± 0.66	14.53 ± 1.06	15.61 ± 0.96	16.49 ± 1.51	17.17 ± 1.40	17.42 ± 2.03	17.74 ± 1.29
Analog FM (SI8-W16-O8)									
PRM (greedy)	19.12 ± 0.83	24.84 ± 0.83	30.80 ± 1.44	34.20 ± 1.53	36.64 ± 1.86	39.52 ± 1.03	41.60 ± 1.56	42.20 ± 2.05	42.84 ± 1.30
PRM (voting)	19.12 ± 0.83	24.84 ± 0.83	30.32 ± 1.26	33.88 ± 0.84	35.48 ± 0.88	37.44 ± 1.16	38.36 ± 0.89	38.72 ± 0.92	38.88 ± 0.53
Voting	19.12 ± 0.83	19.12 ± 0.83	22.88 ± 1.25	29.48 ± 1.05	32.32 ± 0.86	33.92 ± 1.70	34.32 ± 0.33	35.48 ± 0.30	35.52 ± 0.84
Analog FM (SI8-W16 _{hw noise} -O8)									
PRM (greedy)	14.74 ± 0.61	19.87 ± 0.77	25.00 ± 0.39	28.29 ± 1.29	31.69 ± 1.19	33.78 ± 1.46	35.58 ± 1.09	37.42 ± 0.99	39.06 ± 1.61
PRM (voting)	14.74 ± 0.61	19.87 ± 0.77	24.12 ± 0.45	27.05 ± 0.96	30.21 ± 1.32	32.25 ± 1.20	33.01 ± 1.11	33.37 ± 1.02	33.89 ± 0.91
Voting	14.74 ± 0.61	14.74 ± 0.61	18.55 ± 0.69	22.72 ± 0.43	26.84 ± 1.95	28.32 ± 1.22	29.61 ± 1.12	30.05 ± 0.67	30.61 ± 0.81
LLM-QAT (SI8-W4)									
PRM (greedy)	13.40 ± 1.15	17.20 ± 1.00	20.36 ± 1.08	25.40 ± 0.63	29.32 ± 0.35	31.32 ± 1.37	32.80 ± 1.34	33.64 ± 0.82	34.32 ± 0.95
PRM (voting)	13.40 ± 1.15	17.20 ± 1.00	19.56 ± 1.32	23.60 ± 1.11	26.64 ± 0.64	27.88 ± 0.55	29.40 ± 0.92	30.80 ± 1.21	30.64 ± 0.61
Voting	13.40 ± 1.15	13.40 ± 1.15	15.92 ± 1.36	19.68 ± 0.64	21.88 ± 0.95	24.08 ± 0.47	26.44 ± 0.89	27.44 ± 0.57	27.80 ± 0.46
LLM-QAT (SI8-W4 _{hw noise})									
PRM (greedy)	9.36 ± 0.91	11.76 ± 0.98	15.60 ± 1.12	18.52 ± 0.93	22.12 ± 0.83	24.96 ± 0.99	26.36 ± 0.66	27.96 ± 1.28	28.60 ± 0.63
PRM (voting)	9.36 ± 0.91	11.76 ± 0.98	14.80 ± 1.23	17.72 ± 0.95	19.52 ± 0.83	21.04 ± 1.06	21.72 ± 0.79	22.08 ± 0.65	22.28 ± 0.78
Voting	9.36 ± 0.91	9.36 ± 0.91	11.16 ± 1.50	14.32 ± 1.39	15.84 ± 0.56	17.68 ± 0.99	18.68 ± 0.70	19.20 ± 0.79	19.04 ± 1.17

K Hardware throughput and energy efficiency estimation

While deployment on physical AIMC hardware is beyond the scope of this work, we provide estimates of throughput and energy efficiency using an open-source simulator for AIMC-based architectures [26]. We simulated pipelined inference of the *Phi-3-mini-4k-instruct* model with a pre-fill phase of 256 tokens, generation of 64 tokens, and a batch size of 4. Due to the Python-based implementation of the simulator, longer sequences become computationally prohibitive to simulate.

Under these conditions, *Phi-3-mini-4k-instruct* achieves a throughput of 2554.43 tokens/s and an energy efficiency of 199.78 tokens/s/W on the simulated AIMC architecture. To contextualize this

result, we compare against a hypothetical scenario where a GPU could execute all INT8 operations involved in model inference in a single pass at 100% utilization. Even in this idealized comparison, the AIMC-based system demonstrates $3.6\times$ higher energy efficiency.

However, this comparison does not fully capture the potential advantages of AIMC systems. The high density of NVM enables AIMC-based accelerators to host billions of parameters entirely on-chip, eliminating the energy-intensive data movement between off-chip memory and processing units that dominates energy consumption in conventional architectures. This advantage becomes particularly pronounced for models with tens to hundreds of billions of parameters, especially sparse architectures such as MoEs [26]. For such large-scale models, the energy savings from eliminating off-chip memory access can reach multiple orders of magnitude, while also enabling deployment scenarios that would be impractical with conventional hardware due to memory capacity constraints.

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: In the abstract and introduction, we make the following claims. We claim that off-the-shelf LLMs are not able to achieve 4-bit-level performance when deployed on AIMC-based hardware. We demonstrate this in table 1. We claim that our models achieve performance similar to models trained with 4-bit weight and 8-bit static input quantization. This is also demonstrated in table 1. We claim that our models can also be quantized post-training, obtaining competitive performance compared to models trained with state-of-the-art quantization algorithms. This is demonstrated in table 3. We further claim that our models show better scaling behavior under test-time compute scaling than models trained with 4-bit per-channel weight and 8-bit static input quantization. This is shown in figure 4. With respect to related works, we claim that simple STE is sufficient to obtain robustness against static output quantization. This is demonstrated by the results in table 1, and the small drop our model has compared to the off-the-shelf model evaluated in FP16.

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: We discuss the limitations of the work in "Conclusion and Limitations". Our main limitation is the resource intensity of our approach. We acknowledge that this might limit applicability for some researchers and stress the need for more resource-efficient ways of creating analog foundation models. We propose ideas addressing this issue that can be pursued. We also acknowledge the fact that there is still a gap in accuracy that needs to be addressed through further innovation.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.

- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: Our paper is purely empirical.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: In the methods, we describe in detail how we train our analog foundation models. In the appendix, we provide details on the hyperparameters used for training our models and the evaluation we performed, including the exact examples used for few-shot learning, prompt structure, and answer extraction method.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example

- (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
- (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
- (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).
- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [\[Yes\]](#)

Justification: We provide code to reproduce the data-generation and training of our analog foundation models.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyperparameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [\[Yes\]](#)

Justification: All training hyperparameters and evaluation details can be found in the appendix.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [\[Yes\]](#)

Justification: We report error bars for every result involving sampling. Generally, every experiment involving sampling/noise was repeated 10 times and the mean result and standard deviation is reported.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [\[Yes\]](#)

Justification: For the training of our models, we provide detailed resource requirements for training the models, including total training time on different GPUs.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [\[Yes\]](#)

Justification: We made sure to preserve anonymity and conform with the NeurIPS Code of Ethics.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.
- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. **Broader impacts**

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: We discuss that our paper answers an important question in the field of Analog In-Memory Computing, and therefore also in the field of AI accelerators in general. We hope that this inspires new research and motivates further scaling of AIMC-based chips. On a more negative side, we discuss that we based our models on aligned pre-trained models, and that, although we show that our models do not lose the ability to prevent harmful content generation, we still know that these LLMs can produce harmful and toxic content. This is discussed in the Conclusion and Limitations section of the paper.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. **Safeguards**

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [Yes]

Justification: Our models are based on pre-trained models that were aligned post-training. As we show, our models do not lose the ability to prevent generating harmful content or answering to a harmful prompt.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.

- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. Licenses for existing assets

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [Yes]

Justification: We build our models on top of the Phi-3 and Llama 3 series, which we credit. Our dataset is synthetically generated without any starting prompts, so does not warrant credit.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [Yes]

Justification: Code is released under the MIT License. The code is well documented. Apart from the code, no new asset is released.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: No crowdsourcing or research with human subjects was performed.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.
- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: No research with human subjects was performed.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigor, or originality of the research, declaration is not required.

Answer: [NA]

Justification: LLMs were not an important, original, or non-standard component of the core methods in this research.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.