

SQUAT: Stateful Quantization-Aware Training in Recurrent Spiking Neural Networks

Sreyes Venkatesh
ECE, UC Santa Cruz
Santa Cruz, CA, USA
spvenkat@ucsc.edu

Razvan Marinescu
CSE, UC Santa Cruz
Santa Cruz, CA, USA
ramarine@ucsc.edu

Jason K. Eshraghian
ECE, UC Santa Cruz
Santa Cruz, CA, USA
jeshragh@ucsc.edu

Abstract—Weight quantization is used to deploy high-performance deep learning models on resource-limited hardware, enabling the use of low-precision integers for storage and computation. Spiking neural networks (SNNs) share the goal of enhancing efficiency, but adopt an ‘event-driven’ approach to reduce the power consumption of neural network inference. While extensive research has focused on weight quantization, quantization-aware training (QAT), and their application to SNNs, the precision reduction of state variables during training has been largely overlooked, potentially diminishing inference performance. This paper introduces two QAT schemes for stateful neurons: (i) a uniform quantization strategy, an established method for weight quantization, and (ii) threshold-centered quantization, which allocates exponentially more quantization levels near the firing threshold. Our results show that increasing the density of quantization levels around the firing threshold improves accuracy across several benchmark datasets. We provide an ablation analysis of the effects of weight and state quantization, both individually and combined, and how they impact models. Our comprehensive empirical evaluation includes full precision, 8-bit, 4-bit, and 2-bit quantized SNNs, using QAT, stateful QAT (SQUAT), and post-training quantization methods. The findings indicate that the combination of QAT and SQUAT enhance performance the most, but given the choice of one or the other, QAT improves performance by the larger degree. These trends are consistent all datasets. Our methods have been made available in our Python library `snnTorch`: <https://github.com/jeshraghian/snntorch>.

I. INTRODUCTION

The development of low-power neural networks is crucial for enabling operation on portable and edge devices [1]–[3]. Techniques such as pruning [4], specialized data encoding [5], model compression [6], early exiting [7], amongst many others, can be used to reduce the computational cost of running a neural network [8], [9]. In tandem with all these approaches, most edge devices require low or fixed precision model parameters. Quantized neural networks (QNNs) require full precision weights to be approximated down to lower-capacity representations which further reduces memory and computation demands [10], [11].

Spiking neural networks (SNNs), drawing from models of biological neurons, use binarized activations, enabling a neuron to either emit a spike or remain inactive. This binary representation allows neuromorphic processors to bypass certain computations and memory accesses typical in conventional deep learning, leading to significant power savings [12]–[15]. The hidden state of a spiking neuron is often governed by a

dynamical system, and is thought to encode information which is then communicated through the firing pattern of the neuron. These patterns can vary in spike timing, frequency, intervals between spikes, amongst many other theories, offering a range of encoding possibilities [16]–[19].

SNNs are often tailored for edge devices and have demonstrated powerful computational capabilities even with considerably small models [20]–[23], and larger-scale models are also emerging [24]. Fixed-precision representations in SNN accelerators are standard, and quantized SNNs (QSNNs) are commonplace when deploying spike-based models on neuromorphic hardware [25]–[30].

By default, the most widely used deep learning Python libraries train models with full precision parameters. Quantizing a model after it has been trained in full precision is known as ‘Post-Training Quantization’ (PTQ). Performance can be enhanced further by using ‘Quantization-Aware Training’ (QAT), where weights are quantized during the forward-pass but gradients are calculated in full precision. The quantization step is ignored during error backpropagation as it is a non-differentiable function. This process allows for the loss from the model to account for truncation errors made during quantization. This modification to the backpropagation process takes more computational resources, but enhances accuracy.

To improve loss convergence, training QNNs and QSNNs both benefit from modifying the training process [31]–[36]. In general, noise has been regarded as an obstacle towards convergence¹ [40], [41]. Though it was suggested in Ref. [42] that QSNNs are tolerant to truncation, provided that any rounding errors do not trigger a threshold-crossing of the neuron state, thus either hallucinating or eliminating a spike [43].

Most advances in low-precision neural networks apply variations of QAT to learning fixed-precision weights [44]–[46]. At one extreme, Ref. [47] demonstrated promising performance in a ternary language model with over one billion weights. These variants of QAT may involve alternative distributions from which quantization levels are sampled from (e.g., uniform quantization, exponentially distributed quantization). The most common practice is to rely on full precision simulators that emulate quantized weights. This is done by

¹Note that targeted non-systematic noise has been demonstrated to assist in convergence, due to the noise acting as a regularizer [37]–[39], and is not included in this claim

restricting the permitted full precision levels of weights, and then training such models with QAT (e.g., using the Brevitas Python library [48]).

While it may be conceptually trivial to extend this practice to the hidden state of neurons: e.g., quantizing states during the forward-pass of training such that the loss accounts for truncation error of states, it is not done in practice. There are a variety of reasons why this is the case: i) an absence of tools that simplify quantization of states; ii) applying QAT to neuron states is computationally expensive: weights are constant over sequence steps, whereas the state must be re-quantized at every sequence step; iii) modern deep learning is less reliant on stateful and sequential neural networks, and iv) there are typically more weights than there are neurons, so optimizing for weights may be thought of as more important – i.e., it is just ‘easier’ to apply post-quantization to the states once a QSNN is deployed.

This study analyzes two forms of stateful quantization aware training (‘SQUAT’) for training QSNNs: ‘uniform quantization’ and ‘exponential quantization’. Exponential quantization allocates more states about the firing threshold of a neuron as near-threshold activity of a neuron is likely to have a greater impact on downstream neurons. Therefore, higher precision is more useful at this point of criticality.

Beyond proposing SQUAT, we present an empirical analysis evaluating how the quantization of weights and states compare in both the post-quantization and quantization-aware training regimes. A comprehensive evaluation is performed across three different datasets of varying difficulty and modality: an image dataset, FashionMNIST [49], an auditory dataset, Spiking Heidelberg Digits (SHD) [50], and an event-based dataset, the DVS Gesture Dataset [51]. Specifically, we test 2-bit, 4-bit, and 8-bit QSNNs under the following cases:

- PTQ of states and weights, both together and separately, for uniform quantization
- PTQ of states and weights, both together and separately, for exponential quantization
- QAT and SQUAT, both together and separately, for uniform quantization
- QAT and SQUAT, both together and separately, for exponential quantization
- A high precision SNN baseline for each experiment

In total, we conduct 129 experiments across three trials each culminating in insights about the impacts of quantization across weights and states, and how to gain the ‘last mile’ of performance from an SNN. Our findings are summarized below:

- 8-bit QSNNs perform competitively against their full precision counter parts for all tested quantization schemes: i.e., PTQ vs. QAT and Exponential Distribution vs Uniform Distribution.
- 4-bit and 2-bit models are far more sensitive to quantization and require more care during training.
- QAT (weights) and SQUAT (states) combined consistently provided the best performance across all bit-widths

for fixed-precision performance.

- When using either SQUAT or PTQ, exponentially distributed levels centered about the threshold consistently improved accuracy for all experiments.
- QAT-only consistently outperforms SQUAT-only. Given a large computational budget, combining both is most effective. Where there are limited training resources and one technique must be prioritized, QAT is more important than SQUAT.

Finally, we release code to assist other researchers to train QSNNs using SQUAT in a straightforward and intuitive manner.

II. BACKGROUND

A. Spiking Neuron Model

The spiking neuron model adopted is the leaky integrate-and-fire neuron:

$$\tau_m \frac{dU}{dt} = -U + RI, \quad (1)$$

where τ_m is the time constant of membrane potential, U is the membrane potential, R is the passive resistance of the membrane of the neuron, and I is the input current [52]. The above equation when governed by discrete time dynamics and represented in a recurrent manner can be represented by:

$$u_{t+1}^j = \beta u_{t+1}^j + \sum_i w^{ij} z_t^j - z_t^j \theta \quad (2)$$

$$z_t^j = \begin{cases} 1, & \text{if } u_{t+1}^j > \theta \\ 0, & \text{otherwise} \end{cases} \quad (3)$$

where u_{t+1}^j is the membrane potential of neuron j at time t ; β is the inverse time constant of membrane potential; w^{ij} is the synaptic weight between neurons i and j . The neuron is reset by the threshold θ every time a spike z_t^j is emitted. Eq. (2) shows how a spike emitted at each instance of time where the membrane potential exceeds the firing threshold.

B. Hard Thresholds in QSNNs

There are two non-differentiable operations in SNNs: one is from spike generation (Eq. (3)), and another is applied to the weights during QAT and to the states during SQUAT. Addressing these challenges is quite well-established:

- **Spike Non-Differentiability:** The non-differentiability of spikes has been addressed by applying a step function during the forward-pass as per (3), and smoothing it out into a ‘surrogate gradient’ in the backward-pass [53].
- **QAT:** QAT follows a similar approach, though with additional steps as the weight is a learnable parameter:
 - 1) Quantized weights are used for computation during the forward-pass, while the original full precision representations are stored in memory.
 - 2) A loss is calculated based on quantized weights.
 - 3) Gradients are calculated while neglecting the gradient of the quantization operator.

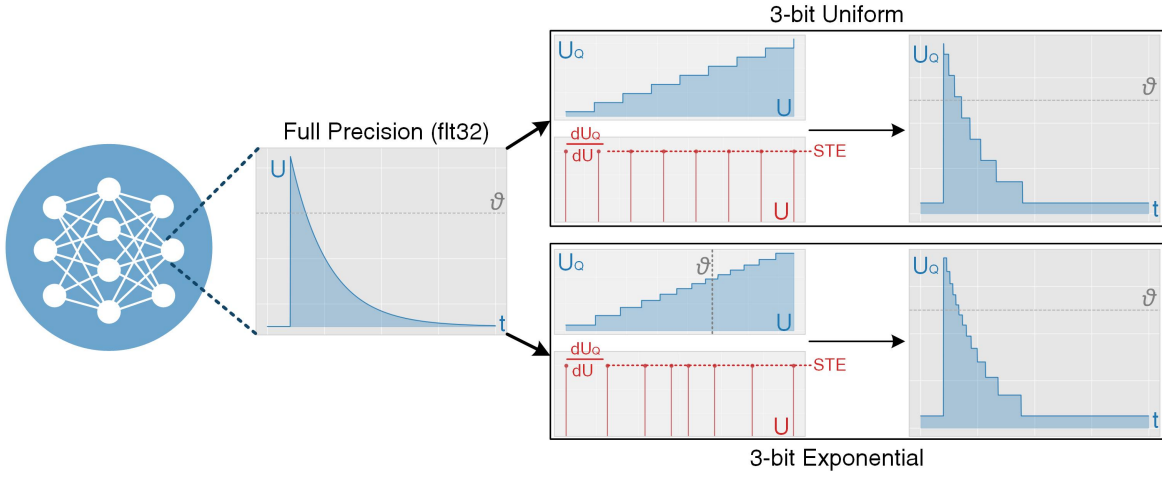


Fig. 1. A graphical depiction of stateful quantization. On the left, a membrane potential trajectory is depicted in full precision. The state can be quantized either via uniform or exponential quantization. A 3-bit (8 levels) quantization scheme is illustrated. In uniform quantization, the permissible levels are evenly distributed. In exponential quantization, the permissible levels are closer about the threshold, and widely distributed moving further away from the threshold. A ‘straight through estimator’ (STE) is also depicted to address the non-differentiability of quantization.

- 4) Weight updates are applied to the full precision weights rather than the quantized weights.
- 5) The process is repeated.

The quantization operator is neglected during the gradient calculation as it is non-differentiable and would otherwise null the gradient signal. To remedy this, a ‘straight through estimator’ (STE) is used to bypass the threshold operator during QAT. The STE acts as an approximate gradient that smooths the thresholding function during training [54]–[58]. More formally, the surrogate gradient of the quantized weight w_q with respect to the real weight w_r is:

$$\frac{\partial w_q}{\partial w_r} = 1 \quad (4)$$

Alternatively, PTQ can be applied to a pre-trained model where weights are quantized after training a full precision model [59]. Whilst computationally cheaper, it also leads to a drop in model performance.

SQUAT extends these principles by calculating a loss that is aware of the quantization of states, while bypassing the non-differentiability using a STE. The next section provides further detail on the variants of SQUAT that we propose and test.

III. METHODS

A. Stateful Quantization-Aware Training

For n -bit quantization, the number of permitted levels Q_l with respect to the number of bits allocated n is $Q_l = 2^n$. We implement SQUAT by distributing these Q_l levels using two different methods, both of which are illustrated in Fig. 1:

- **Uniform Quantization:** The permissible quantization levels are uniformly distributed between the minimum $U_{\min}$ and maximum $U_{\max}$ membrane potentials. $U_{\min}$ and $U_{\max}$ are recalculated for each forward-pass.
- **Exponential Quantization:** Similar to uniform quantization, the maximum and minimum possible voltages

for the membrane potential are applied as the upper and lower bound of permissible levels. However, exponentially more quantization levels are allocated for the threshold. In deep learning, exponentially distributed weights are applied about ‘0’. However, spiking neurons only communicate with one another when the membrane potential reaches the threshold. Intuitively, it stands to reason that a neuron requires more precision about the threshold as this is the regime where network activity is determined.

1) *Uniform SQUAT:* More formally, in uniform quantization, we divide the range between the minimum ($U_{\min}$) and maximum ($U_{\max}$) membrane potentials into equal intervals. Let’s denote the number of intervals as $N = 2^n - 1$. The quantized value U_q can be calculated as follows:

$$U_q = U_{\min} + \left\lceil \frac{U - U_{\min}}{\Delta U} \right\rceil \cdot \Delta U$$

where $\Delta U = \frac{U_{\max} - U_{\min}}{N}$ is the size of each interval. The operator $\lceil \cdot \rceil$ rounds the argument to the nearest integer. It is assumed that $U_{\min}$ and $U_{\max}$ have been scaled and offset such that $U_{\min} = 0$ and $U_{\max} = N$.

2) *Exponential SQUAT:* In this case, the quantization becomes finer as it approaches the threshold U_{th} . The quantization is described using a pair of exponential functions, one for values below the threshold and another for values above it, while being clipped at $U_{\min}$ and $U_{\max}$. The quantized value U_q can be expressed as:

$$U_q = \begin{cases} U_{\min} + \Delta U \cdot \left\lceil \frac{1 - e^{-a(U - U_{\min})}}{\Delta U} \right\rceil, & \text{if } U < \theta \\ U_{\max} - \Delta U \cdot \left\lceil \frac{1 - e^{-b(U_{\max} - U)}}{\Delta U} \right\rceil, & \text{if } U \geq \theta \end{cases}$$

where a and b are the exponents that determine the steepness of the exponential curves below and above the threshold, respec-

tively. This formulation allows for exponential quantization of the membrane potential while considering the resolution implied by the number of bits N .

3) *Straight-Through-Estimator*: As with the STE from (4), the same is applied to the membrane potential:

$$\frac{\partial U_q}{\partial U} = 1 \quad (5)$$

These equations provide a mathematical framework for quantizing the membrane potential in both uniform and exponentially varying manners near the threshold.

B. Testing

We assess performance across three data sets: FashionMNIST, Spiking Heidelberg, and DVS Gesture. For each dataset:

- A full precision baseline is obtained on a lightweight architecture.
- A PTQ baseline is obtained by converting the full precision baseline to 8-bits, 4-bits, and 2-bits for weights-only, states-only, and both. For the case where only the weights are quantized, the states are left in full-precision and vice versa.
- QAT-only results are obtained by retraining the model across 8-bit, 4-bit, and 2-bit weights. States are left in full precision.
- SQUAT-only results are obtained by retraining the model across 8-bit, 4-bit, and 2-bit weights, for uniform quantization and exponential quantization. Weights are left in full precision.
- QAT and SQUAT are both applied across 8-bit, 4-bit, and 2-bit weights, for uniform quantization and exponential quantization

C. Training

For all experiments, a threshold-shifted arc-tangent surrogate gradient is used to deal with the spike non-differentiability:

$$\frac{\delta z}{\delta U} = \frac{1}{\pi} \frac{1}{[1 + (\pi U \alpha)^2]} \quad (6)$$

Each network is tested with a cosine annealing learning rate scheduler and the Adam Optimizer [60], [61]. To maintain the hardware benefits obtained from operating on lower-bit variables, we intentionally expose the network to a discontinuous loss landscape with flat surfaces and use the more challenging task of using the total spike count per neuron as the logits. Many state-of-the-art results opt to use a ‘read-out’ layer instead, where a standard artificial neuron node is applied at the final layer instead of a spiking layer.

For the DVS-Gesture dataset we use the Mean-Squared Error loss with respect to the spikes of each output neuron: z_t^j and the target spike count c^j . Each of those losses is then summed over M output classes.

$$\mathcal{L}_{MSE} = \sum_j^M \sum_t (c^j - z_t^j)^2, \quad (7)$$

For the FashionMNIST and Spiking Heidelberg Digits datasets we use the cross entropy loss as applied to the spike count:

$$\mathcal{L}_{CE} = \frac{1}{C} \sum_j^M \sum_{t=0}^C N(\log(p^j[t]), Y^j) \quad (8)$$

where C is the number of time steps, $p^j[t]$ is the softmax probability of the spike count of the output neuron j at time step t , N is the negative log likelihood loss function, and Y^j is the target spike count of output neuron j .

D. Model Architecture

Given the notation $C_{\text{out}}\text{Conv}k$ and $N_{\text{in}}\text{Dense}N_{\text{out}}$, where C_{out} is the number of output channels, k is the kernel dimension, and N_{in} and N_{out} are the input-output dimensions of a linear layer, the model architectures used were:

- FashionMNIST: 16Conv5-MP2-64Conv5-MP2-1024Dense10. Batch normalization is applied to all convolutional layers [62].
- DVS Gesture: 16Conv5-MP2-32Conv5-MP2-8800Dense11. Dropout is applied during training. Batch normalization is applied to all convolutional layers.
- Spiking Heidelberg Digits: 700Dense1000-1000Dense20. Each linear layer is followed by batch normalization operation. Dropout is also applied.

IV. EXPERIMENTAL RESULTS

Brevitas was used to apply uniformly quantize the network weight and bias parameters during training and testing; snnTorch was used to instantiate the SNNs, and PyTorch 2.0 for training and testing. Each result shown is the average taken across three trials.

The FashionMNIST and SHD dataset simulations were limited to 100 epochs with early stopping of 20 epochs applied. The DVS Gesture Dataset accuracy was obtained by running 500 epochs. More details are given in the individual sections of the datasets.

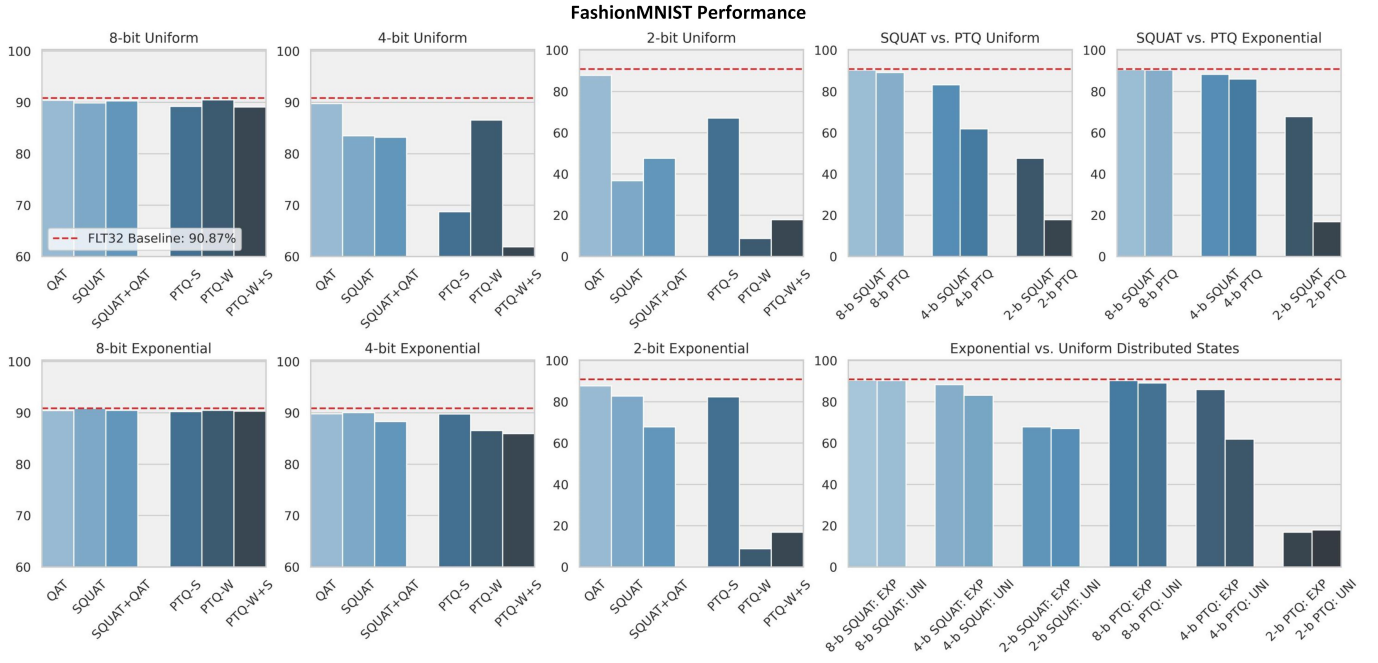
TABLE I
FULL PRECISION ACCURACY

Dataset	Accuracy
FMNIST	90.87
SHD	79.78
DVS	86.24

A. FashionMnist

The FashionMNIST dataset contains ten classes of clothing items and accessories [49]. The raw FashionMNIST dataset was repeatedly passed to the network for 25 time steps of simulation without encoding.

Table II displays the results of exponentially-distributed states, Table III displays the results of uniformly-distributed states, and Fig. 2 shows barplots of the same results organized to provide easier visual comparisons. Quantizing weights to



8-bits had negligible impact on accuracy, regardless of the quantization configuration (SQUAT/QAT or PTQ). At 4-bits, exponentially-distributed states using SQUAT significantly outperforms uniformly-distributed states with SQUAT (Table III) in all test cases by a significant margin. Finally, 2-bit weights destroy performance when using PTQ, but the model can be salvaged by combining QAT and SQUAT using exponentially distributed levels centered about the threshold (i.e., an accuracy boost from 16.81% to 67.82%).

Note PTQ and QAT of weights are the same between uniform and threshold-centered quantization, as quantization levels is only applied to states, and these two configurations do not quantize the states. The same is true for all datasets, but are included for ease of reference.

TABLE II
FMNIST EXPONENTIALLY CENTERED QUANTIZED PERFORMANCE

	8-bit	4-bit	2-bit
QAT States	90.87	90.04	82.79
QAT Weights	90.43	89.79	87.77
QAT Weights and States	90.49	88.29	67.82
PTQ States	90.26	89.78	82.36
PTQ Weights	90.48	86.56	8.72
PTQ Weights and States	90.32	85.95	16.81

TABLE III
FMNIST UNIFORM QUANTIZED PERFORMANCE

	8-bit	4-bit	2-bit
QAT States	89.91	83.50	36.76
QAT Weights	90.43	89.79	87.77
QAT Weights and States	90.31	83.21	47.58
PTQ States	89.22	68.72	67.06
PTQ Weights	90.48	86.56	8.72
PTQ Weights and States	89.11	61.87	17.87

B. Spiking Heidelberg Digits

The Spiking Heidelberg Digits is an audio based classification dataset of 20 possible output classes. There are approximately 10,000 recordings of spoken digits from 0 to 9 in both English and German [50].

The SHD dataset was passed to the network for 25 time steps during training and 30 time steps during testing. Table IV displays the average test accuracy over 100 epochs across 3 trials each for exponentially-distributed states, Table V displays the same data for uniformly-distributed states, and Fig. 3 illustrates the same data as a barplot.

Similarly to the FashionMNIST dataset, 8-bit quantization had little to no effect on accuracy, regardless of the quantization configuration. At 4-bits, exponentially-distributed quantization of states improved performance relative to PTQ. At 2-bits, there is significant variation across all configurations. At this point, when both states and weights are quantized down

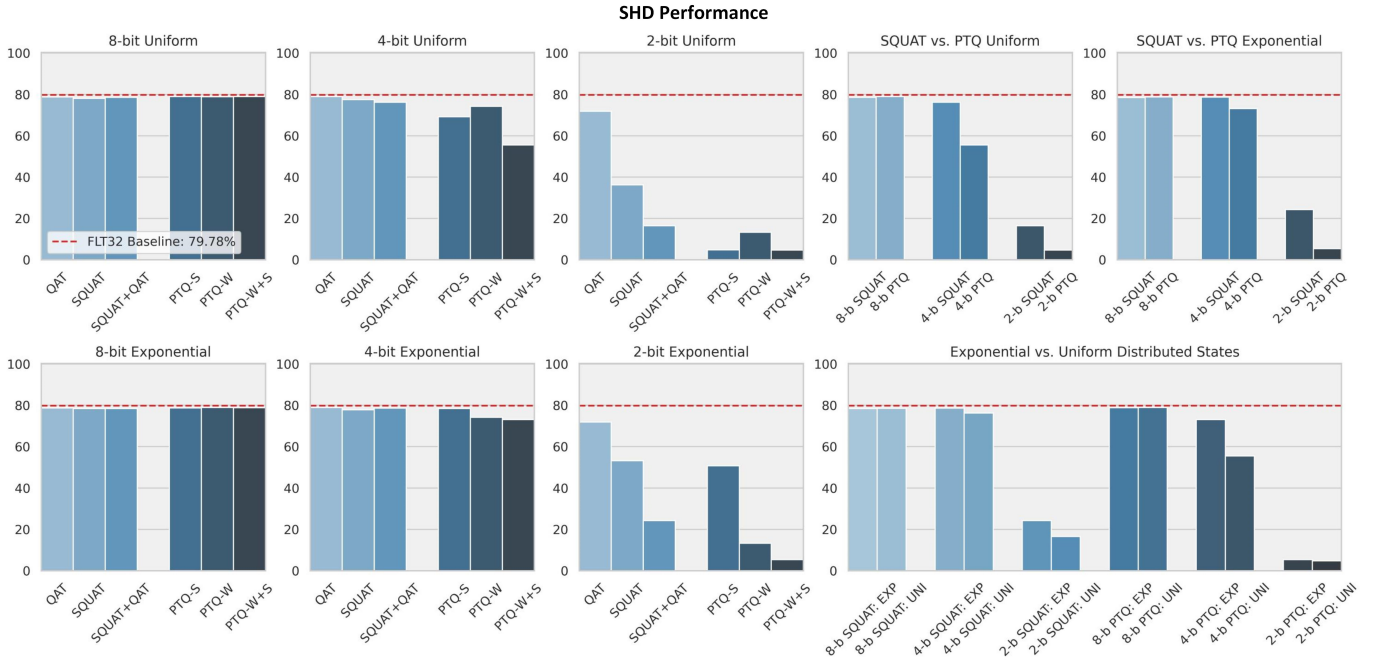


Fig. 3. SHD performance. Top row: (i) 8/4/2-b uniformly distributed states across QAT (N-b weights, flt32 states), SQUAT (flt32 weights, N-b states), SQUAT+QAT (N-b weights, N-b states), and PTQ-S (flt32 weights, N-b states), PTQ-W (N-b weights, flt32 states), PTQ-W+S (N-b weights, N-b states). (ii) SQUAT vs PTQ uniformly and exponentially distributed states (SQUAT and QAT are both applied across N-b weights and states, and compared against PTQ of N-b states and weights). Bottom row: (iii) 8/4/2-b exponentially distributed states across QAT (N-b weights, flt32 states), SQUAT (flt32 weights, N-b states), SQUAT+QAT (N-b weights, N-b states), and PTQ-S (flt32 weights, N-b states), PTQ-W (N-b weights, flt32 states), PTQ-W+S (N-b weights, N-b states). (iv) Comparison between exponential and uniformly distributed states: SQUAT+QAT are used across N-b states and weights, then N-b PTQ is used across N-b states and weights.

to 2-bits, these models are quite unstable. But using QAT and SQUAT together nonetheless outperforms PTQ methods by a significant margin. Exponentially distributed state levels outperforms uniformly distributed levels by approximately 8% at 2-bits.

TABLE IV
SHD EXPONENTIALLY CENTERED QUANTIZED PERFORMANCES

	8-bit	4-bit	2-bit
QAT States	78.48	77.81	53.25
QAT Weights	78.72	78.98	71.88
QAT Weights and States	78.46	78.63	24.26
PTQ States	78.76	78.42	50.70
PTQ Weights	78.89	74.19	13.27
PTQ Weights and States	78.79	73.10	5.33

TABLE V
SHD UNIFORM CENTERED QUANTIZED PERFORMANCE

	8-bit	4-bit	2-bit
QAT States	78.03	77.43	36.27
QAT Weights	78.72	78.98	71.88
QAT Weights and States	78.52	76.23	16.48
PTQ States	78.96	69.16	4.74
PTQ Weights	78.89	74.19	13.27
PTQ Weights and States	78.94	55.47	4.69

C. DVS Gesture

The DVS Gesture dataset is an event based dataset with 11 output classes consisting of various hand gestures [51]. The DVS dataset was passed to the network for 25 time steps during training and 150 time steps during testing. Of all the three datasets, the DVS gesture dataset was the most sensitive to quantization. This may be because it has the largest number of neurons and synapses in its architecture, and thus, the most truncation error accumulated in the loss. With reference to Table VI for exponentially distributed quantizations, Table VII for uniformly distributed levels, and Fig. 4 for the corresponding barplots, all PTQ training runs at 2-bits yielded rather hopeless models despite our best hyperparameter sweep efforts. However, the combination of exponentially-distributed levels, together with QAT and SQUAT boosted performance back up to 79.89%, as against 9.59% using PTQ (i.e., random chance), and 28.03% when using uniform distributed SQUAT + QAT. This highlights the significant importance of adopting better distribution strategies in the extreme quantization regime. The trends for all three datasets remain incredibly consistent, with huge performance margins for exponentially-distributed hidden state levels in the extreme quantization regime.

V. DISCUSSION

Deducing the optimal quantization scheme: For the 8-bit experiments one can see that the performance of the model

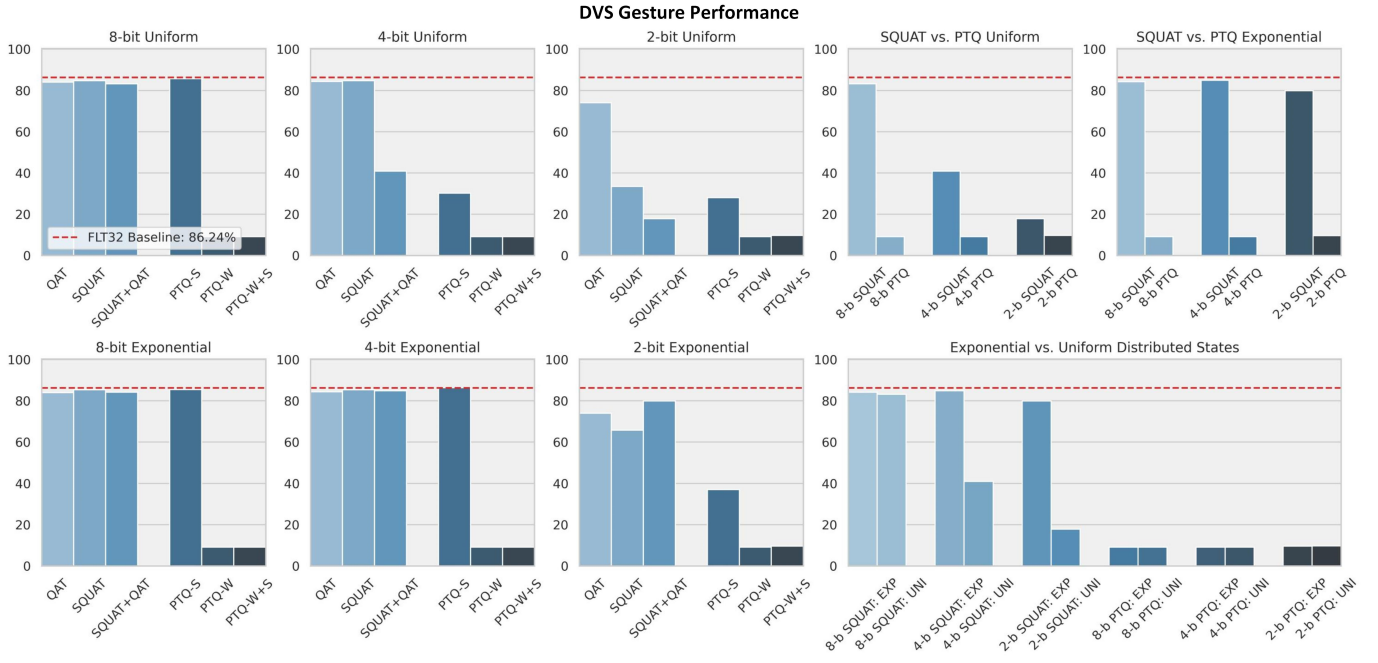


Fig. 4. DVS Gesture Dataset performance. Top row: (i) 8/4/2-b uniformly distributed states across QAT (N-b weights, flt32 states), SQUAT (flt32 weights, N-b states), SQUAT+QAT (N-b weights, N-b states), and PTQ-S (flt32 weights, N-b states), PTQ-W (N-b weights, flt32 states), PTQ-W+S (N-b weights, N-b states). (ii) SQUAT vs PTQ uniformly and exponentially distributed states (SQUAT and QAT are both applied across N-b weights and states, and compared against PTQ of N-b states and weights). Bottom row: (iii) 8/4/2-b exponentially distributed states across QAT (N-b weights, flt32 states), SQUAT (flt32 weights, N-b states), SQUAT+QAT (N-b weights, N-b states), and PTQ-S (flt32 weights, N-b states), PTQ-W (N-b weights, flt32 states), PTQ-W+S (N-b weights, N-b states). (iv) Comparison between exponential and uniformly distributed states: SQUAT+QAT are used across N-b states and weights, then N-b PTQ is used across N-b states and weights.

TABLE VI
DVS EXPONENTIALLY CENTERED QUANTIZED PERFORMANCE

	8-bit	4-bit	2-bit
QAT States	85.35	85.35	65.78
QAT Weights	83.97	84.34	73.99
QAT Weights and States	84.22	84.85	79.89
PTQ States	85.48	86.36	36.95
PTQ Weights	9.09	9.09	9.09
PTQ Weights and States	9.09	9.09	9.59

TABLE VII
DVS UNIFORM CENTERED QUANTIZED PERFORMANCE

	8-bit	4-bit	2-bit
QAT States	84.72	84.72	33.46
QAT Weights	83.97	84.34	73.99
QAT Weights and States	83.20	40.91	17.80
PTQ States	85.73	30.17	28.03
PTQ Weights	9.09	9.09	9.09
PTQ Weights and States	9.09	9.09	9.72

was independent of the given quantization scheme, which is generally quite unsurprising [63]. Whereas for the 4-bit and 2-bit scenarios, it becomes increasingly clear that the quantization method (i.e., PTQ/QAT/SQUAT/exponential/uniform-distribution) is a major factor on the performance of the network. Overall, performing QAT while also implementing exponentially-distributed quantization about the threshold on the hidden states of the neurons results in the least degradation

of accuracy across low-bit quantization schemata. This is in line with our other findings of exponentially distributed levels outperforming uniformly distributed states, along with QAT performing better than PTQ.

Membrane potentials were reaching much higher values than previously anticipated: This strongly affected the considerations that had to be taken when optimizing the quantization of membrane potentials. A specific example being that: it was expected that the reset mechanism would have a strong influence on what the membrane potential is after it exceeds the threshold. In practice, the optimal thresholds tended to be between $0.5 < \theta < 1$, whereas the membrane potential was reaching values as high as 20. Based on (2), it is clear that the reset mechanism was having a negligible effect on the membrane potential. Analyzing the membrane potential traces meant an alternative approach to clipping was required.

Clipping had a significant impact on the accuracy: Quantization sets maximum and minimum bounds on the membrane potential. Originally, it was believed that since a spike is determined by the value of membrane potential with respect to the threshold, it would be unnecessary to store membrane potentials that exceed approximately twice the threshold based on (1). However, we found that clipping must encompass the maximum and minimum membrane potential values that would be reached during the full precision training

process to optimize performance. This implies that allowing membrane potentials to reach larger values allows the network to better mitigate noise. Further investigation may allocate less levels to negative hidden states, and act more closely to rectification units as these values do not contribute to spiking activity.

Threshold centering can mitigate a significant loss of accuracy at lower bits: For 8-bit quantization there was no significant difference as to whether the quantization was uniform across the range of membrane potentials or whether the quantization allocated more bits for values centered around the threshold. This implies that the difference in membrane potentials between time steps was larger than the smallest difference between the allocated state-levels. However for 4-bit quantization, threshold centering performed nearly as well as 8-bit across all datasets, whereas for 4-bit uniform quantization of states, accuracy becomes significantly worse as compared to 8-bit uniform quantization. The variation of accuracy for 4-bit uniform quantization was also much larger than the variation of accuracy for 8-bit uniform quantization.

Future Asymmetric Quantizations: One can see that neuron states are more robust to quantization during PTQ whereas network weights are more adaptable to quantization during QAT. This aligns with the knowledge that noise from the quantization would be somewhat absorbed by the threshold dynamics of the neuron. Whereas it is clear that noise brought in to the weights after the network is finished training often leads the algorithm to poor performance. Allocating additional memory on certain parts of the network, for example, 3-bits for weights and 5- bits for states with QAT or 6-bits for weights and 2-bits for states, would allow the user to maximize the networks capacity for quantization and performance.

VI. CONCLUSION

We have demonstrated the necessity of adopting both QAT and SQUAT in enabling extreme-quantization regimes of QSNNs to maintain some semblance of reasonable performance. Our findings clearly indicate that threshold centering of exponentially distributed states is better at mitigating the performance degradation that comes with lower precision. We have also demonstrated the differences in performance for QAT and PQT occur at lower bit (ie. 4-bit and 2-bit) situations. Enabling stateful quantization-aware training is effectively as simple as passing a single argument to a leaky integrate-and-fire neuron in snnTorch, simplifying the overall process of obtaining additional accuracy. Both uniform and exponential quantization schemes are available and paramterizable. Code snippets demonstrating how to use it are provided.

```

1 import torch
2 import snntorch as snn
3 from snntorch.functional import quant
4
5 # neuron parameters
6 beta = 0.5
7 thr = 5
8
9 # random input
10 rand_input = torch.rand(1)
11
12 # uniform quantization
13 q_uni = quant.state_quant(num_bits=4,
14     ↪ uniform=True, threshold=thr)
15 lif_1 = snn.Leaky(beta=beta,
16     ↪ threshold=thr, state_quant=q_uni)
17
18 # forward-pass for one step
19 spk, mem = lif(rand_input, mem)
20
21 # exponential quantization
22 q_exp = quant.state_quant(num_bits=4,
23     ↪ uniform=False, threshold=thr)
24 lif_2 = snn.Leaky(beta=beta,
25     ↪ threshold=thr, state_quant=q_exp)

```

REFERENCES

- [1] Hande Alemdar, Vincent Leroy, Adrien Prost-Boucle, and Frédéric Pétrot. Ternary neural networks for resource-efficient ai applications. In *2017 International Joint Conference on Neural Networks (IJCNN)*, pages 2547–2554, 2017.
- [2] Eunhyeok Park, Dongyoung Kim, Soobom Kim, Yong-Deok Kim, Gunhee Kim, Sungroh Yoon, and Sungjoo Yoo. Big/little deep neural network for ultra low power inference. In *2015 International Conference on Hardware/Software Codesign and System Synthesis (CODES+ISSS)*, pages 124–132, 2015.
- [3] Okan Kopuklu, Neslihan Kose, Ahmet Gunduz, and Gerhard Rigoll. Resource efficient 3d convolutional neural networks. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) Workshops*, Oct 2019.
- [4] Steven A Janowsky. Pruning versus clipping in neural networks. *Physical Review A*, 39(12):6600, 1989.
- [5] Daniel Auge, Julian Hille, Etienne Mueller, and Alois Knoll. A survey of encoding techniques for signal processing in spiking neural networks. *Neural Processing Letters*, 53(6):4693–4710, 2021.
- [6] Yu Cheng, Duo Wang, Pan Zhou, and Tao Zhang. Model compression and acceleration for deep neural networks: The principles, progress, and challenges. *IEEE Signal Processing Magazine*, 35(1):126–136, 2018.
- [7] Surat Teerapittayanon, Bradley McDanel, and Hsiang-Tsung Kung. Branchynet: Fast inference via early exiting from deep neural networks. In *2016 23rd international conference on pattern recognition (ICPR)*, pages 2464–2469. IEEE, 2016.
- [8] Song Han, Huizi Mao, and William J Dally. Deep compression: Compressing deep neural networks with pruning, trained quantization and huffman coding. *arXiv preprint arXiv:1510.00149*, 2015.
- [9] Yihui He, Xiangyu Zhang, and Jian Sun. Channel pruning for accelerating very deep neural networks. In *Proceedings of the IEEE International Conference on Computer Vision (ICCV)*, Oct 2017.
- [10] Itay Hubara, Matthieu Courbariaux, Daniel Soudry, Ran El-Yaniv, and Yoshua Bengio. Quantized neural networks: Training neural networks with low precision weights and activations, 2016.
- [11] Shu-Chang Zhou, Yu-Zhi Wang, He Wen, Qin-Yao He, and Yu-Heng Zou. Balanced quantization: An effective and efficient approach to quantized neural networks. *Journal of Computer Science and Technology*, 32:667–682, 2017.
- [12] Michael Pfeiffer and Thomas Pfeil. Deep learning with spiking neurons: Opportunities and challenges. *Frontiers in neuroscience*, 12:774, 2018.

- [13] Mostafa Rahimi Azghadi, Corey Lammie, Jason K Eshraghian, Melika Payvand, Elisa Donati, Bernabe Linares-Barranco, and Giacomo Indiveri. Hardware implementation of deep network accelerators towards healthcare and biomedical applications. *IEEE Transactions on Biomedical Circuits and Systems*, 14(6):1138–1159, 2020.
- [14] Kaushik Roy, Akhilesh Jaiswal, and Priyadarshini Panda. Towards spike-based machine intelligence with neuromorphic computing. *Nature*, 575(7784):607–617, 2019.
- [15] Samuel Schmidgall, Jascha Achterberg, Thomas Miconi, Louis Kirsch, Rojin Ziaei, S Hajiseydrizi, and Jason Eshraghian. Brain-inspired learning in artificial neural networks: a review. *arXiv preprint arXiv:2305.11252*, 2023.
- [16] Balint Petro, Nikola Kasabov, and Rita M Kiss. Selection and optimization of temporal spike encoding methods for spiking neural networks. *IEEE transactions on neural networks and learning systems*, 31(2):358–370, 2019.
- [17] Shimul Kanti Nath, Sujan Kumar Das, Sanjoy Kumar Nandi, Chen Xi, Camilo Verbel Marquez, Armando Rúa, Mutsunori Uenuma, Zhongrui Wang, Songqing Zhang, Rui-Jie Zhu, et al. Optically tunable electrical oscillations in oxide-based memristors for neuromorphic computing. *Advanced Materials*, page 2400904, 2024.
- [18] Kwabena Boahen. Dendrocentric learning for synthetic intelligence. *Nature*, 612(7938):43–50, 2022.
- [19] Jason Kamran Eshraghian, Kyoungrok Cho, Ciyan Zheng, Minho Nam, Herbert Ho-Ching Iu, Wen Lei, and Kamran Eshraghian. Neuromorphic vision hybrid rram-cmos architecture. *IEEE Transactions on Very Large Scale Integration (VLSI) Systems*, 26(12):2816–2829, 2018.
- [20] Sumit B Shrestha and Garrick Orchard. Slayer: Spike layer error reassignment in time. *Advances in neural information processing systems*, 31, 2018.
- [21] Alexander Henkes, Jason K Eshraghian, and Henning Wessels. Spiking neural network for nonlinear regression. *arXiv preprint arXiv:2210.03515*, 2022.
- [22] Yikai Yang, Jason K Eshraghian, Nhan Duy Truong, Armin Nikpour, and Omid Kavehei. Neuromorphic deep spiking neural networks for seizure detection. *Neuromorphic Computing and Engineering*, 3(1):014010, 2023.
- [23] Sami Barchid, José Mennesson, Jason Eshraghian, Chaabane Djéraba, and Mohammed Bennamoun. Spiking neural networks for frame-based and event-based single object localization. *Neurocomputing*, 559:126805, 2023.
- [24] Rui-Jie Zhu, Qihang Zhao, Guoqi Li, and Jason K Eshraghian. Spikept: Generative pre-trained language model with spiking neural networks. *arXiv preprint arXiv:2302.13939*, 2023.
- [25] Mike Davies, Narayan Srinivasa, Tsung-Han Lin, Gautham China, Yongqiang Cao, Sri Harsha Choday, Georgios Dimou, Prasad Joshi, Nabil Imam, Shweta Jain, et al. Loihi: A neuromorphic manycore processor with on-chip learning. *Ieee Micro*, 38(1):82–99, 2018.
- [26] Garrick Orchard, E Paxon Frady, Daniel Ben Dayan Rubin, Sophia Sanborn, Sumit Bam Shrestha, Friedrich T Sommer, and Mike Davies. Efficient neuromorphic signal processing with loihi 2. In *2021 IEEE Workshop on Signal Processing Systems (SiPS)*, pages 254–259. IEEE, 2021.
- [27] Hannah Bos and Dylan Muir. Sub-mw neuromorphic snn audio processing applications with rockpool and xylo. *Embedded Artificial Intelligence: Devices, Embedded Systems, and Industrial Applications*, page 69, 2023.
- [28] Ole Richter, Yannan Xing, Michele De Marchi, Carsten Nielsen, Mercurios Katsimpris, Roberto Cattaneo, Yudi Ren, Qian Liu, Sadique Sheik, Tugba Demirci, et al. Speck: A smart event-based vision sensor with a low latency 327k neuron convolutional neuronal network processing pipeline. *arXiv preprint arXiv:2304.06793*, 2023.
- [29] Jens E Pedersen, Steven Abreu, Matthias Jobst, Gregor Lenz, Vittorio Fra, Felix C Bauer, Dylan R Muir, Peng Zhou, Bernhard Vogginger, Kade Heckel, et al. Neuromorphic intermediate representation: a unified instruction set for interoperable brain-inspired computing. *arXiv preprint arXiv:2311.14641*, 2023.
- [30] Mostafa Rahimi Azghadi, Ying-Chen Chen, Jason K Eshraghian, Jia Chen, Chih-Yang Lin, Amirali Amirsalemani, Adnan Mehonic, Anthony J Kenyon, Burt Fowler, Jack C Lee, et al. Complementary metal-oxide semiconductor and memristive hardware for neuromorphic computing. *Advanced Intelligent Systems*, 2(5):1900189, 2020.
- [31] Mike Davies, Narayan Srinivasa, Tsung-Han Lin, Gautham China, Yongqiang Cao, Sri Harsha Choday, Georgios Dimou, Prasad Joshi, Nabil Imam, Shweta Jain, Yuyun Liao, Chit-Kwan Lin, Andrew Lines, Ruokun Liu, Deepak Mathaikutty, Steven McCoy, Arnab Paul, Jonathan Tse, Guruguhathan Venkataramanan, Yi-Hsin Weng, Andreas Wild, Yoonseok Yang, and Hong Wang. Loihi: A neuromorphic manycore processor with on-chip learning. *IEEE Micro*, 38(1):82–99, 2018.
- [32] Amar Shrestha, Haowen Fang, Daniel Patrick Rider, Zaidao Mei, and Qinru Qiu. In-hardware learning of multilayer spiking neural networks on a neuromorphic processor. In *2021 58th ACM/IEEE Design Automation Conference (DAC)*, pages 367–372, 2021.
- [33] Emre O Neftci, Charles Augustine, Somnath Paul, and Georgios Detorakis. Event-driven random back-propagation: Enabling neuromorphic deep learning machines. *Frontiers in neuroscience*, 11:324, 2017.
- [34] Amar Shrestha, Haowen Fang, Qing Wu, and Qinru Qiu. Approximating back-propagation for a biologically plausible local learning rule in spiking neural networks. In *Proceedings of the International Conference on Neuromorphic Systems*, pages 1–8, 2019.
- [35] Steve K Esser, Rathinakumar Appuswamy, Paul Merolla, John V Arthur, and Dharmendra S Modha. Backpropagation for energy-efficient neuromorphic computing. *Advances in neural information processing systems*, 28, 2015.
- [36] Peter U. Diehl and Matthew Cook. Efficient implementation of stdp rules on spinnaker neuromorphic hardware. In *2014 International Joint Conference on Neural Networks (IJCNN)*, pages 4288–4295, 2014.
- [37] Zhonghui You, Jinmian Ye, Kunming Li, Zenglin Xu, and Ping Wang. Adversarial noise layer: Regularize neural network by adding noise. In *2019 IEEE International Conference on Image Processing (ICIP)*, pages 909–913, 2019.
- [38] Kam-Chuen Jim, C.L. Giles, and B.G. Horne. An analysis of noise in recurrent neural networks: convergence and generalization. *IEEE Transactions on Neural Networks*, 7(6):1424–1438, 1996.
- [39] Hyeonwoo Noh, Tackgeun You, Jonghwan Mun, and Bohyung Han. Regularizing deep neural networks by noise: Its interpretation and optimization. *Advances in neural information processing systems*, 30, 2017.
- [40] Gunhan Dundar and Kenneth Rose. The effects of quantization on multilayer neural networks. *IEEE Transactions on Neural Networks*, 6(6):1446–1451, 1995.
- [41] H. Djahanshahi, M. Ahmadi, G.A. Jullien, and W.C. Miller. Quantization noise improvement in a hybrid distributed-neuron ann architecture. *IEEE Transactions on Circuits and Systems II: Analog and Digital Signal Processing*, 48(9):842–846, 2001.
- [42] Jason K. Eshraghian, Corey Lammie, Mostafa Rahimi Azghadi, and Wei D. Lu. Navigating local minima in quantized spiking neural networks, 2022.
- [43] Jason K Eshraghian, Xinxin Wang, and Wei D Lu. Memristor-based binarized spiking neural networks: Challenges and applications. *IEEE Nanotechnology Magazine*, 16(2):14–23, 2022.
- [44] Nitin Rathi, Priyadarshini Panda, and Kaushik Roy. Stdp-based pruning of connections and weight quantization in spiking neural networks for energy-efficient recognition. *IEEE Transactions on Computer-Aided Design of Integrated Circuits and Systems*, 38(4):668–677, 2019.
- [45] Hin Wai Lui and Emre Neftci. Hessian aware quantization of spiking neural networks. In *International Conference on Neuromorphic Systems 2021*, pages 1–5, 2021.
- [46] Sayeed Shafayet Chowdhury, Isha Garg, and Kaushik Roy. Spatio-temporal pruning and quantization for low-latency spiking neural networks, 2021.
- [47] Shuming Ma, Hongyu Wang, Lingxiao Ma, Lei Wang, Wenhui Wang, Shaohan Huang, Li Dong, Ruiping Wang, Jilong Xue, and Furu Wei. The era of 1-bit llms: All large language models are in 1.58 bits. *arXiv preprint arXiv:2402.17764*, 2024.
- [48] Alessandro Pappalardo. Xilinx/brevitas, 2023.
- [49] Han Xiao, Kashif Rasul, and Roland Vollgraf. Fashion-mnist: a novel image dataset for benchmarking machine learning algorithms, 2017.
- [50] Benjamin Cramer, Yannik Stradmann, Johannes Schemmel, and Friedemann Zenke. The heidelberg spiking data sets for the systematic evaluation of spiking neural networks. *IEEE Transactions on Neural Networks and Learning Systems*, 33(7):2744–2757, 2022.
- [51] Arnon Amir, Brian Taba, David Berg, Timothy Melano, Jeffrey McK-instry, Carmelo Di Nolfo, Tapan Nayak, Alexander Andreopoulos, Guillaume Garreau, Marcela Mendoza, Jeff Kusnitz, Michael Debole, Steve Esser, Tobi Delbruck, Myron Flickner, and Dharmendra Modha. A low power, fully event-based gesture recognition system. In *2017*

IEEE Conference on Computer Vision and Pattern Recognition (CVPR), pages 7388–7397, 2017.

- [52] Eric Hunsberger and Chris Eliasmith. Spiking deep networks with lif neurons, 2015.
- [53] Emre O. Neftci, Hesham Mostafa, and Friedemann Zenke. Surrogate gradient learning in spiking neural networks, 2019.
- [54] Yoshua Bengio, Nicholas Léonard, and Aaron Courville. Estimating or propagating gradients through stochastic neurons for conditional computation, 2013.
- [55] Penghang Yin, Jiancheng Lyu, Shuai Zhang, Stanley Osher, Yingyong Qi, and Jack Xin. Understanding straight-through estimator in training activation quantized neural nets, 2019.
- [56] Alexander Shekhovtsov and Viktor Yanush. Reintroducing straight-through estimators as principled methods for stochastic binary networks, 2021.
- [57] Zechun Liu, Kwang-Ting Cheng, Dong Huang, Eric Xing, and Zhiqiang Shen. Nonuniform-to-uniform quantization: Towards accurate quantization via generalized straight-through estimation, 2022.
- [58] Ting-Han Fan, Ta-Chung Chi, Alexander I. Rudnicky, and Peter J. Ramadge. Training discrete deep generative models via gapped straight-through estimator, 2022.
- [59] Rachmad Vidya Wicaksana Putra and Muhammad Shafique. Q-SpiNN: A framework for quantizing spiking neural networks. In *2021 International Joint Conference on Neural Networks (IJCNN)*. IEEE, jul 2021.
- [60] Tong He, Zhi Zhang, Hang Zhang, Zhongyue Zhang, Junyuan Xie, and Mu Li. Bag of tricks for image classification with convolutional neural networks, 2018.
- [61] Ilya Loshchilov and Frank Hutter. Sgdr: Stochastic gradient descent with warm restarts, 2017.
- [62] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. Dropout: a simple way to prevent neural networks from overfitting. *The journal of machine learning research*, 15(1):1929–1958, 2014.
- [63] Yury Nahshan, Brian Chmiel, Chaim Baskin, Evgenii Zheltonozhskii, Ron Banner, Alex M Bronstein, and Avi Mendelson. Loss aware post-training quantization. *Machine Learning*, 110(11-12):3245–3262, 2021.