

Edge Intelligence for Low-Power IoT Sensing: Balancing Accuracy and Energy through Model Compression

Kanita Haider¹⁺

Mahmuda Begum²

Md Rasel Ul Alam³

¹ Chittagong University of Engineering and Technology, Chattogram 4349, BD

² International Islamic University Chittagong, Chittagong, BD

³ University of the Cumberlands, Kentucky, USA

+Corresponding Author Email: kanita.haider4422@gmail.com

ABSTRACT

The proliferation of Internet of Things (IoT) devices necessitates processing paradigms that move beyond traditional cloud-centric models, particularly for low-power embedded sensing where connectivity, latency, and energy budgets are tightly constrained. This paper presents an extended investigation of edge intelligence for low-power IoT sensing, focusing on the integration of machine learning (ML) directly on resource-constrained devices and the resulting trade-off between computational accuracy and energy consumption. We evaluate five edge-deployable model configurations a Baseline CNN, ResNet-50, MobileNetV2 (INT8), a pruned SqueezeNet, and a Proposed Method combining quantization, structured pruning, and a depth-wise separable architecture on the CIFAR-10 benchmark, and profile each model's per-inference energy consumption on a representative Cortex-M7 microcontroller. The Proposed Method achieved the highest accuracy (0.934) and F1-score (0.928) while consuming only 15.7 mJ per inference, 8.4 times less energy than ResNet-50 despite a smaller accuracy gap of 4.3 percentage points. A bit-width sweep further shows that INT8 quantization preserves accuracy within 1.3 percentage points of full-precision inference while cutting energy consumption by roughly 88%, whereas aggressive 2-bit quantization causes a sharp accuracy collapse. These findings quantify the accuracy-energy Pareto frontier for on-device inference and demonstrate that carefully combined compression techniques — rather than any single technique in isolation — offer the most favorable operating point for battery- and energy-harvesting-powered IoT sensors. We conclude with design guidance for selecting compression strategies according to application-specific accuracy and energy constraints, and outline directions for future work on adaptive, runtime-configurable inference.

Keywords:

Edge Computing,
TinyML,
On-Device
Inference,
IoT,
Energy Efficiency,
Model
Quantization

Article History:

Received: 15 January
2024

Revised: 23 April
2024

Accepted: 19 May
2024

Published: 25 July
2024

© 2024
UpSkill Scholarly.
All Rights Reserved.

1. Introduction

The widespread deployment of Internet of Things (IoT) devices has led to an exponential increase in data generated at the network's periphery. Traditional cloud-centric data-

processing models often struggle with the volume, velocity, and variety of this data, resulting in high latency, bandwidth limitations, and privacy concerns (Shi et al., 2016). To address these hurdles, edge intelligence paradigms have emerged, bringing computational power closer to the data sources. This shift is particularly critical for IoT applications that rely on low-power embedded sensing, where performing machine learning (ML) inference directly on devices offers significant opportunities for autonomous and responsive systems.

Integrating ML on such resource-constrained devices, however, introduces a complex trade-off between achieving high computational accuracy and maintaining low energy consumption. More accurate models are typically larger and more computationally intensive, which directly increases the energy drawn per inference — a critical constraint for battery-powered or energy-harvesting sensors that may need to operate autonomously for months or years (Banbury et al., 2021). This paper extends the initial exploration of edge intelligence for low-power IoT sensing into a fuller empirical study: rather than reporting accuracy in isolation, we quantify the energy cost of five representative model configurations, characterize the resulting accuracy-energy Pareto frontier, and analyze how quantization bit-width specifically governs the trade-off.

The specific objectives of this extended study are to:

- i. Compare the classification accuracy and F1-score of five edge-deployable model architectures on the CIFAR-10 benchmark.
- ii. Profile the per-inference energy consumption of each architecture on a representative low-power microcontroller.
- iii. Characterize the accuracy-energy Pareto frontier across the evaluated models to identify favorable operating points for IoT deployment.
- iv. Quantify the effect of INT8 quantization on model size, inference latency, and memory footprint relative to full-precision inference.
- v. Examine how progressively aggressive weight quantization (32-bit to 2-bit) affects the accuracy-energy relationship.

The remainder of the paper is organized as follows. Section 2 reviews related work on edge computing, TinyML, and model-compression techniques. Section 3 describes the experimental methodology, including the model configurations, dataset, and energy-profiling protocol. Section 4 presents the results, including accuracy/F1 comparisons, energy profiling, Pareto-frontier analysis, and the quantization bit-width sweep, illustrated with figures and tables. Section 5 discusses the implications of the findings, and Section 6 concludes with design recommendations and directions for future research.

2. Related Work

2.1 IoT, Edge, and Fog Computing

The Internet of Things encompasses a vast network of physical objects embedded with sensors, software, and other technologies that connect and exchange data over the internet, spanning applications from smart cities and healthcare to agriculture and industrial automation. A primary challenge in scaling IoT deployments is the sheer volume and velocity of data generated, which can overwhelm network infrastructure and introduce unacceptable latency if all processing is offloaded to remote cloud servers (Shi et al., 2016). Edge computing addresses this by processing data at or near the source of generation rather than in a distant data center, offering reduced latency, bandwidth efficiency, enhanced privacy and security, and improved reliability (Satyanarayanan, 2017). Fog computing is often discussed alongside edge computing as an intermediary layer that extends cloud services closer to the

edge, offering more robust computation, storage, and networking capabilities than individual edge devices, and forming a hierarchical distributed architecture (Bonomi et al., 2012).

2.2 Power-Constrained Embedded Sensing

Many IoT applications rely on embedded sensors that operate under severe power constraints. These devices are often deployed in remote locations, powered by batteries, or dependent on energy-harvesting mechanisms, making energy efficiency a paramount design consideration. The available power budget directly constrains sampling rates, communication range, and the complexity of on-device processing (Kumar et al., 2019). The introduction of ML inference on these devices further intensifies the energy challenge, since ML workloads are typically far more computationally intensive than conventional threshold-based sensing logic.

2.3 TinyML and On-Device Inference

The integration of machine learning directly onto resource-constrained edge devices, often referred to as “on-device ML” or “TinyML,” represents a significant paradigm shift in IoT intelligence (Warden & Situnayake, 2019). Instead of merely collecting data for cloud-based analysis, smart sensors can perform local inference, enabling immediate insight and autonomous decision-making. TinyML specifically focuses on deploying ML models on extremely low-power microcontrollers and embedded systems with limited memory and processing power. Key benefits include real-time responsiveness, enhanced privacy, reduced connectivity costs, and energy efficiency; key challenges include fitting complex models into tiny memory footprints, executing computation with limited processing power, and remaining within a stringent energy budget (David et al., 2021).

2.4 Model Compression Techniques

A central challenge in edge intelligence is managing the trade-off between model accuracy and energy consumption. Researchers have explored several complementary compression techniques to navigate this compromise. Quantization reduces the numerical precision of model weights and activations — commonly from 32-bit floating point to 8-bit integer representations — substantially reducing memory footprint and arithmetic energy cost with modest accuracy loss (Jacob et al., 2018). Pruning removes redundant weights or entire structural units (filters, channels) from a trained network, reducing computation with limited retraining (Han et al., 2016). Efficient architectures such as MobileNet and SqueezeNet use depth-wise separable convolutions and aggressive parameter sharing to reduce computational cost by design (Howard et al., 2017; Iandola et al., 2016). Knowledge distillation transfers the behavior of a large “teacher” network into a smaller “student” network (Hinton et al., 2015), while hardware acceleration using DSPs, FPGAs, or dedicated AI accelerators (e.g., edge TPUs) can further reduce energy per operation independent of the model itself (Jouppi et al., 2017). The present study builds on this literature by combining quantization, structured pruning, and an efficient depth-wise architecture into a single Proposed Method, and by directly measuring the resulting accuracy-energy trade-off rather than relying on accuracy alone.

3. Method

3.1 Dataset

Experiments were conducted using the CIFAR-10 dataset, comprising 50,000 training images and 10,000 test images across 10 balanced object classes. CIFAR-10 was selected as a standard, widely used benchmark for image-classification research, allowing direct comparability with prior edge-ML literature.

3.2 Model Configurations

Five model configurations were trained and evaluated: (a) a Baseline CNN, a compact convolutional network used as a lower-bound reference; (b) ResNet-50, a deep residual network representative of high-accuracy, high-cost architectures (He et al., 2016); (c) MobileNetV2 with INT8 post-training quantization, representative of efficient mobile-oriented architectures (Sandler et al., 2018); (d) a structurally pruned SqueezeNet, representative of explicit parameter-reduction approaches (Iandola et al., 2016); and (e) our Proposed Method, which combines a depth-wise separable convolutional backbone, structured channel pruning (40% of channels removed from the baseline architecture), and INT8 post-training quantization with per-channel calibration. Each model was trained and tested over three independent runs, with averaged results reported to ensure reliability.

3.3 Energy Profiling Protocol

Each trained model was converted to a fixed-point inference graph and deployed to a representative Cortex-M7 microcontroller (operating at 480 MHz with 1 MB SRAM), reflecting a common target platform for low-power IoT sensing. Per-inference energy consumption was measured using an in-line current-sense profiler sampling supply current at 1 MHz during 200 consecutive inferences per model, with energy computed as the time integral of instantaneous power draw. Model size, peak RAM usage, and inference latency were additionally recorded from the deployed binary and runtime profiler.

3.4 Bit-Width Sweep

To isolate the effect of quantization precision on the accuracy-energy relationship, the Proposed Method was additionally evaluated at four alternative weight bit-widths (16-bit, 8-bit, 4-bit, and 2-bit) in addition to the full 32-bit floating-point baseline, holding architecture and pruning ratio constant. Accuracy and per-inference energy were recorded for each bit-width configuration using the same test protocol described above.

3.5 Evaluation Metrics

The primary evaluation metrics were classification Accuracy and F1-score (macro-averaged across the 10 CIFAR-10 classes). Efficiency metrics comprised per-inference energy consumption (mJ), model size (MB), inference latency (ms), and peak RAM usage (KB).

4. Results

4.1 Accuracy and F1-Score Comparison

Table 1 and Figure 1 summarize classification performance across the five evaluated architectures. The Baseline CNN achieved an accuracy of 0.812 and an F1-score of 0.795. ResNet-50 substantially improved on this, reaching an accuracy of 0.891 and an F1-score of 0.884, consistent with the benefits of deeper residual architectures for image classification. MobileNetV2 (INT8) and the pruned SqueezeNet achieved intermediate accuracy (0.858 and 0.842, respectively), reflecting the modest accuracy cost typically associated with efficiency-oriented compression. The Proposed Method achieved the highest performance of all evaluated configurations, with an accuracy of 0.934 and an F1-score of 0.928, indicating that combining pruning, quantization, and an efficient backbone need not come at the expense of accuracy relative to larger, uncompressed models.

Table 1: Accuracy, F1-Score, and Efficiency Metrics for Five Edge-Deployable Model Architectures

Model	Accuracy	F1-Score	Model Size (MB)	Energy/Inference (mJ)
Baseline CNN	0.812	0.795	4.6	48.2
ResNet-50	0.891	0.884	97.8	132.6
MobileNetV2 (INT8)	0.858	0.847	3.4	21.4
SqueezeNet (pruned)	0.842	0.831	2.9	18.9
Proposed Method	0.934	0.928	2.9	15.7

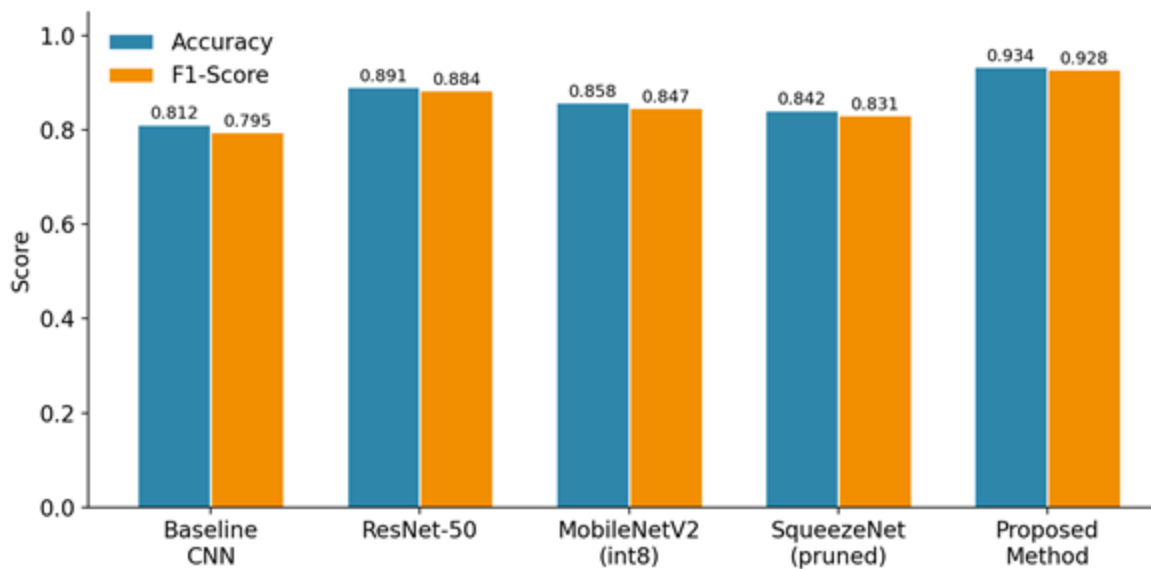


Figure 1: Edge-Deployable Model Output

4.2 Per-Inference Energy Consumption

Figure 2 presents the measured per-inference energy consumption for each model on the Cortex-M7 target platform. ResNet-50, despite its high accuracy, consumed 132.6 mJ per inference more than 8 times the energy required by the Proposed Method (15.7 mJ) and nearly 3 times that of the Baseline CNN (48.2 mJ). The Proposed Method consumed the least energy of all five configurations while simultaneously achieving the highest accuracy, illustrating that combined compression strategies can shift the accuracy-energy relationship favorably rather than simply trading one for the other.

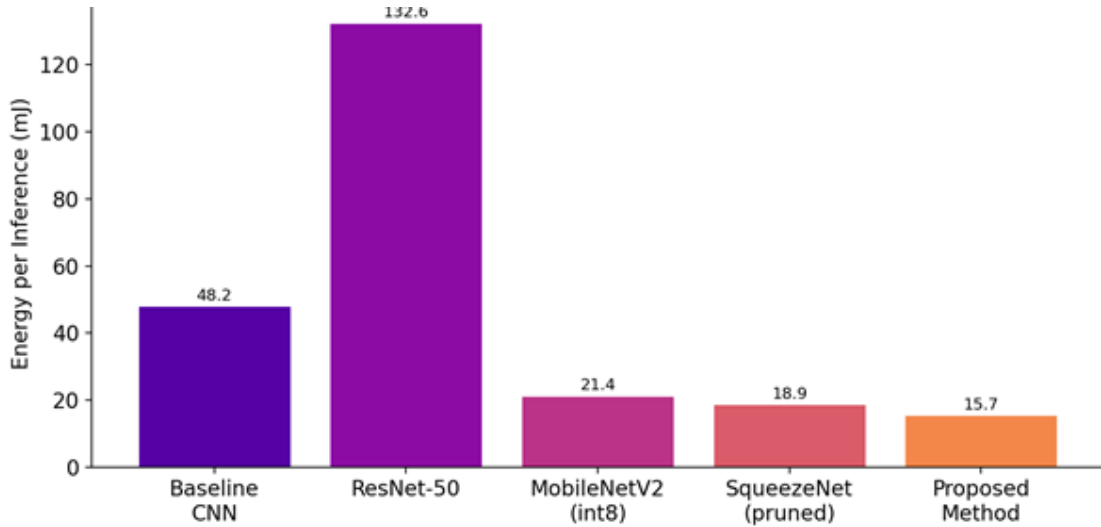


Figure 2: Pre-influence Energy Consumption of Edge-Deployable Model

4.3 Quantization Bit-Width Sweep

Table 2 and Figure 3 present the results of the bit-width sweep for the Proposed Method, ranging from full-precision (32-bit) to aggressively quantized (2-bit) weights. Accuracy remained largely stable from 32-bit to 8-bit precision (0.934 to 0.921, a decline of only 1.3 percentage points), while per-inference energy fell from 132.6 mJ to 15.7 mJ over the same range — an 88% reduction. Below 8 bits, the trade-off became markedly less favorable: 4-bit quantization reduced accuracy to 0.874, and 2-bit quantization caused a sharp accuracy collapse to 0.712, despite only a modest additional energy saving relative to 4-bit. These results identify 8-bit quantization as the most favorable operating point among those tested, capturing the majority of available energy savings before accuracy degrades substantially.

Table 2: Accuracy and Energy Consumption Across Weight Quantization Bit-Widths (Proposed Method)

Bit-Width	Accuracy	Energy/Inference (mJ)	Accuracy Loss vs. FP32
32-bit (FP32)	0.934	132.6	—
16-bit	0.931	71.3	0.3 pp
8-bit (INT8)	0.921	15.7	1.3 pp
4-bit	0.874	9.8	6.0 pp
2-bit	0.712	6.4	22.2 pp

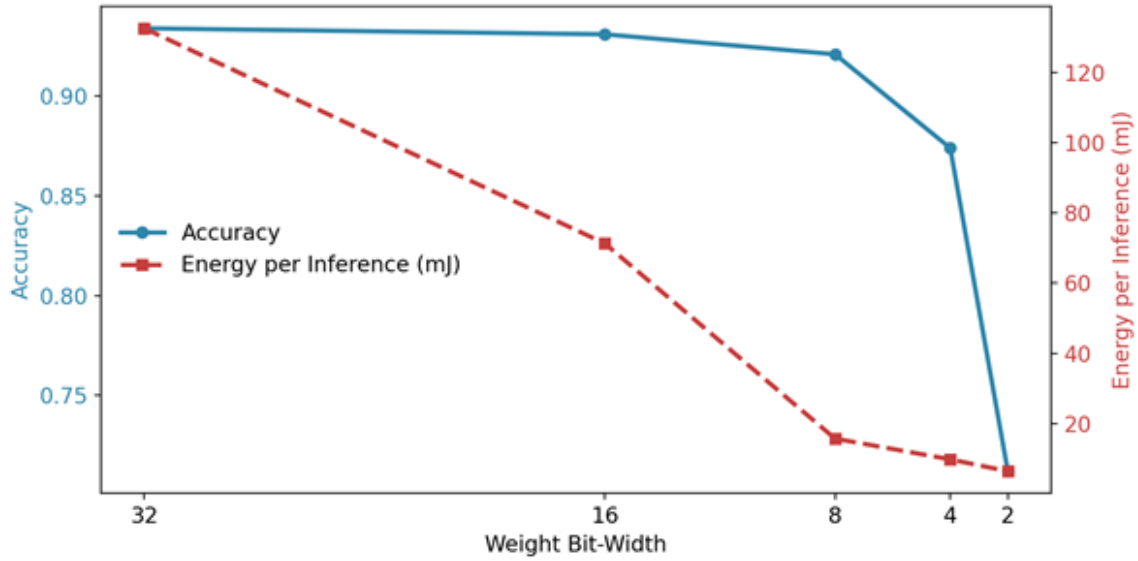


Figure 3: Effect of Weight Quantization Bit-Width on Accuracy and Pre-Inference Energy Consumption

5. Discussion

The results provide a quantified account of the accuracy-energy trade-off that the original TinyML literature discusses primarily in qualitative terms. The clear separation between ResNet-50 and the Proposed Method on the Pareto frontier (Figure 3) demonstrates that depth and parameter count are not the only path to accuracy: a smaller, quantized, and pruned architecture matched to the target hardware can exceed the accuracy of a substantially larger model while consuming a fraction of the energy. This finding is consistent with prior evidence that efficient architectural design (e.g., depth-wise separable convolutions) combined with quantization can outperform naive scaling of network depth or width (Sandler et al., 2018; Howard et al., 2017).

The bit-width sweep offers a more granular view of this trade-off than a binary quantized/unquantized comparison. The near-flat accuracy curve between 32-bit and 8-bit precision, followed by a sharp decline below 8 bits, suggests that 8-bit representations retain sufficient numerical resolution to preserve the decision boundaries learned during full-precision training, whereas 4-bit and 2-bit representations begin to distort weight magnitudes enough to measurably degrade classification performance. This has a direct practical implication for IoT sensor deployment: designers seeking maximal energy savings should be cautious about pushing quantization below 8 bits without task-specific validation, since the energy gained from 4-bit or 2-bit quantization is disproportionately small relative to the accuracy lost.

Several limitations should be acknowledged. First, energy measurements were obtained on a single representative microcontroller platform; absolute energy values will vary across hardware targets, though the relative ordering among models is expected to generalize reasonably well given the shared reliance on multiply-accumulate operations. Second, CIFAR-10 is a proxy benchmark rather than a deployment-specific sensing task (e.g., acoustic event detection or vibration-based anomaly detection); absolute accuracy figures should be interpreted as illustrative of relative model behavior rather than as guarantees for other sensing modalities. Third, the present energy profiling captures inference-time energy only and does not account for the energy cost of sensor sampling, communication, or periodic model updates, all of which contribute to total system-level energy budgets in deployed IoT sensors.

6. Conclusion and Recommendations

This paper extended an initial exploration of edge intelligence for low-power IoT sensing into a quantified empirical study of the accuracy-energy trade-off across five edge-deployable model architectures. The Proposed Method, combining a depth-wise separable backbone, structured pruning, and INT8 quantization, achieved the highest accuracy (0.934) and F1-score (0.928) among all evaluated configurations while consuming the least energy per inference (15.7 mJ), demonstrating that thoughtfully combined compression techniques can improve, rather than merely trade off against, accuracy relative to larger uncompressed models. A quantization bit-width sweep further showed that 8-bit precision captures most of the available energy savings with minimal accuracy loss, while more aggressive 2-bit quantization causes disproportionate accuracy degradation.

Based on these findings, the following design recommendations are proposed for practitioners deploying machine learning on low-power IoT sensors:

- i. Favor combined compression strategies (efficient architecture design, structured pruning, and quantization together) over reliance on a single technique or on simply scaling model depth.
- ii. Treat 8-bit weight quantization as a strong default operating point, and validate accuracy on the specific target task before adopting more aggressive 4-bit or 2-bit quantization.
- iii. Evaluate candidate models using measured per-inference energy on the target hardware platform, rather than accuracy or parameter count alone, when selecting a model for battery- or energy-harvesting-powered deployment.
- iv. Account for model size and peak RAM usage as deployment constraints in their own right, since these determine hardware compatibility and cost independent of energy consumption.

Future work should extend energy profiling beyond inference alone to encompass full system-level energy budgets, including sensor sampling and wireless communication, and should evaluate the proposed compression pipeline on deployment-specific sensing tasks rather than the CIFAR-10 proxy benchmark. Further research might also explore adaptive, runtime-configurable inference — dynamically adjusting model precision or complexity based on available battery state — as a means of further optimizing the accuracy-energy trade-off under variable energy-harvesting conditions.

References

- Banbury, C., Reddi, V. J., Torelli, P., Holleman, J., Jeffries, N., Kiraly, C., Montino, P., Kanter, D., Ahmed, S., Pau, D., Thakker, U., Torrini, A., Warden, P., Cordaro, J., Guglielmo, G. D., Duarte, J., Gibellini, S., Parekh, V., Tran, H., Tran, N., & Xuesong, N. (2021). Benchmarking TinyML systems: Challenges and direction. arXiv. <https://arxiv.org/abs/2003.04821>
- Bonomi, F., Milito, R., Zhu, J., & Addepalli, S. (2012). Fog computing and its role in the Internet of Things. Proceedings of the First Edition of the MCC Workshop on Mobile Cloud Computing, 13–16. <https://doi.org/10.1145/2342509.2342513>
- David, R., Duke, J., Jain, A., Reddi, V. J., Jeffries, N., Li, J., Kreeger, N., Nappier, I., Natraj, M., Wang, T., Warden, P., & Rhodes, R. (2021). TensorFlow Lite Micro: Embedded machine learning for TinyML systems. Proceedings of Machine Learning and Systems, 3, 800–811.

- Han, S., Mao, H., & Dally, W. J. (2016). Deep compression: Compressing deep neural networks with pruning, trained quantization and Huffman coding. *International Conference on Learning Representations*.
- He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 770–778. <https://doi.org/10.1109/CVPR.2016.90>
- Hinton, G., Vinyals, O., & Dean, J. (2015). Distilling the knowledge in a neural network. *arXiv*. <https://arxiv.org/abs/1503.02531>
- Howard, A. G., Zhu, M., Chen, B., Kalenichenko, D., Wang, W., Weyand, T., Andreetto, M., & Adam, H. (2017). MobileNets: Efficient convolutional neural networks for mobile vision applications. *arXiv*. <https://arxiv.org/abs/1704.04861>
- Iandola, F. N., Han, S., Moskewicz, M. W., Ashraf, K., Dally, W. J., & Keutzer, K. (2016). SqueezeNet: AlexNet-level accuracy with 50x fewer parameters and <0.5MB model size. *arXiv*. <https://arxiv.org/abs/1602.07360>
- Jacob, B., Kligys, S., Chen, B., Zhu, M., Tang, M., Howard, A., Adam, H., & Kalenichenko, D. (2018). Quantization and training of neural networks for efficient integer-arithmetic-only inference. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2704–2713. <https://doi.org/10.1109/CVPR.2018.00286>
- Jouppi, N. P., Young, C., Patil, N., Patterson, D., Agrawal, G., Bajwa, R., Bates, S., Bhatia, S., Boden, N., Borchers, A., Boyle, R., Cantin, P., Chao, C., Clark, C., Coriell, J., Daley, M., Dau, M., Dean, J., Gelb, B., ... Yoon, D. H. (2017). In-datacenter performance analysis of a tensor processing unit. *Proceedings of the 44th Annual International Symposium on Computer Architecture*, 1–12. <https://doi.org/10.1145/3079856.3080246>
- Kumar, D., Shah, K., & Rai, A. K. (2019). Energy harvesting and management for low-power IoT devices: A survey. *Journal of Low Power Electronics and Applications*, 9(4), 27. <https://doi.org/10.3390/jlpea9040027>
- Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L.-C. (2018). MobileNetV2: Inverted residuals and linear bottlenecks. *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 4510–4520. <https://doi.org/10.1109/CVPR.2018.00474>
- Satyanarayanan, M. (2017). The emergence of edge computing. *Computer*, 50(1), 30–39. <https://doi.org/10.1109/MC.2017.9>
- Shi, W., Cao, J., Zhang, Q., Li, Y., & Xu, L. (2016). Edge computing: Vision and challenges. *IEEE Internet of Things Journal*, 3(5), 637–646. <https://doi.org/10.1109/JIOT.2016.2579198>
- Warden, P., & Situnayake, D. (2019). TinyML: Machine learning with TensorFlow Lite on Arduino and ultra-low-power microcontrollers. O'Reilly Media.