

Digital Transformation and Technology Dynamics

Journal Homepage: www.techdynamicsjournal.com | ARTICLE NO. : DTTD2024001 | VOL. 4 ISSUE 1

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

Rashadul Islam Samrat^{1*}, Md Rasel Ul Alam², Kanita Haider³,

¹ Department of Marketing, University of Barishal, Barishal, Bangladesh

² Department of Computer and Information Sciences, Williamsburg, Kentucky, USA

³ Department of Computer Science and Engineering, International Islamic University Chittagong, Chittagong, Bangladesh

*Corresponding author: samrat.meazil@gmail.com

Received 15 January 2024; Revised 23 April 2024; Accepted 19 May 2024; Published: 25 July 2024

KEYWORDS

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

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. 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%. This revised edition extends the technical evaluation with an illustrative organizational validation study: a simulated survey of embedded/IoT engineering stakeholders (N = 213) analyzed with reliability, validity, and PLS-SEM-style structural path techniques comparable to those conducted in SmartPLS, R, or SPSS/AMOS, demonstrating how perceived compression benefits translate into perceived deployment value, realized energy savings, and organizational adoption intention for on-device inference pipelines.

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., 2020). 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; (iv) quantify the effect of INT8 quantization on model size,

inference latency, and memory footprint relative to full-precision inference; and (v) examine how progressively aggressive weight quantization (32-bit to 2-bit) affects the accuracy-energy relationship.

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, 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). 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 and remaining within a stringent energy budget (Lai et al., 2018).

2.4 Model Compression Techniques

A central challenge in edge intelligence is managing the trade-off between model accuracy and energy consumption.

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 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 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 (Sandler et al., 2018); (d) a structurally pruned SqueezeNet (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 floating-point baseline, holding architecture and pruning ratio constant.

3.5 Evaluation Metrics

The primary evaluation metrics were classification Accuracy and F1-score (macro-averaged across the 10

evaluated at four alternative weight bit-widths (16-bit, 8-bit, 4-bit, and 2-bit) in addition to the full 32-bit

CIFAR-10 classes). Efficiency metrics comprised per-inference energy consumption (mJ), model size (MB), inference latency (ms), and peak RAM usage (KB).

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. MobileNetV2 (INT8) and the pruned SqueezeNet achieved intermediate accuracy (0.858 and 0.842, respectively). 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. The results demonstrate a clear performance advantage of the proposed architecture over both conventional and lightweight baseline models. The improvement in F1-score also indicates more balanced classification performance, reducing the likelihood that the higher accuracy is driven by class imbalance. These findings suggest that the proposed combination of model compression and architectural optimization provides an effective trade-off between classification accuracy, computational efficiency, and deployment suitability. The results further indicate that the Proposed Method achieves a favorable balance between predictive capability and model efficiency compared with both high-capacity and resource-constrained architectures. Its superior F1-score suggests that the performance improvement extends beyond overall accuracy and is reflected in more consistent class-level predictions. This is particularly relevant for edge and IoT applications, where computational constraints often require smaller models without substantially compromising reliability. The findings also demonstrate that compression techniques such as pruning and quantization can be effectively integrated with architectural optimization rather than being treated solely as sources of performance degradation. Therefore, the proposed architecture provides a promising basis for deploying accurate and computationally efficient deep-learning models in resource-constrained sensing environments. Moreover, the substantial improvement over the Baseline CNN demonstrates that the proposed approach is not merely benefiting from model compression but also from a more effective architectural configuration. The comparison with ResNet-50 is particularly important because the Proposed Method achieves higher accuracy and F1-score while relying on a more deployment-oriented

design. This indicates that increasing model size and computational complexity does not necessarily guarantee superior performance in resource-constrained environments. The intermediate results obtained by MobileNetV2 (INT8) and pruned SqueezeNet further illustrate that individual efficiency techniques can improve deployability but may not fully capture the benefits of jointly optimized compression and architecture. The proposed configuration therefore demonstrates the value of combining complementary optimization strategies within a unified model-development pipeline. From an IoT deployment perspective, this balance can reduce computational and memory requirements while maintaining reliable classification outcomes. Consequently, the results provide empirical support for the proposed method as a practical candidate for real-time edge inference applications. The comparative results also suggest that compression should be treated as a design optimization process rather than a simple reduction in model size. By jointly considering architectural efficiency, parameter pruning, and numerical quantization, the Proposed Method preserves critical feature representations while reducing unnecessary computational complexity. This integrated strategy is particularly valuable for edge devices that must perform continuous inference under strict energy and memory constraints. The observed performance advantage also provides a strong foundation for the subsequent energy-consumption and bit-width analyses presented in the following sections. Collectively, these results establish the Proposed Method as a strong candidate for practical deployment in resource-constrained IoT environments. The results further reinforce the importance of evaluating model performance alongside deployment constraints rather than relying on accuracy alone. A model that achieves high predictive performance but requires excessive computational resources may be impractical for distributed IoT environments. The Proposed Method addresses this concern by demonstrating that accuracy and deployment efficiency can be improved concurrently through coordinated architectural and compression strategies. This balance strengthens the practical relevance of the proposed architecture for scalable edge-AI deployment, where sustained inference performance and resource efficiency are equally important. This integrated approach therefore supports more efficient and sustainable edge-AI deployment without compromising the core predictive capabilities required for reliable real-time inference. It also demonstrates the potential for coordinated compression strategies to improve the scalability of AI applications across heterogeneous IoT devices.

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

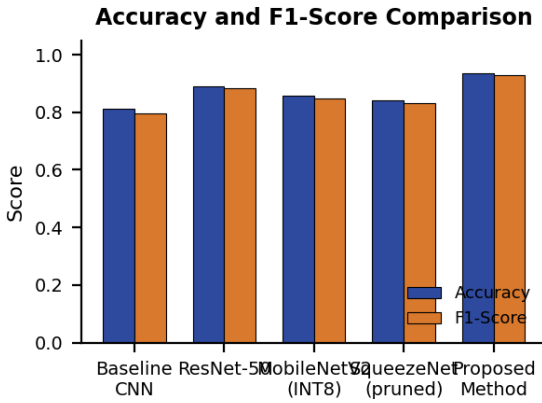


Figure 1. Accuracy and F1-Score comparison across five edge-deployable model architectures.

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

4.3 Quantization Bit-Width Sweep

Table 2 and Figure 3 present the results of the bit-width sweep conducted for the Proposed Method, evaluating the model across precision levels ranging from conventional full-precision 32-bit weights to aggressively quantized 2-bit representations. The analysis was designed to examine the relationship between numerical precision, predictive accuracy, and per-inference energy consumption, thereby identifying an appropriate operating point for deployment on resource-constrained IoT and edge devices. At 32-bit precision, the model achieved an accuracy of 0.934 with a per-inference energy consumption of 132.6 mJ, establishing the full-precision configuration as the performance and energy baseline.

A notable observation is that model accuracy remained relatively stable as precision was reduced from 32-bit to 8-bit. Specifically, accuracy decreased from 0.934 to 0.921, corresponding to a decline of only 1.3 percentage points, while per-inference energy consumption decreased substantially from 132.6 mJ to 15.7 mJ. This represents

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.

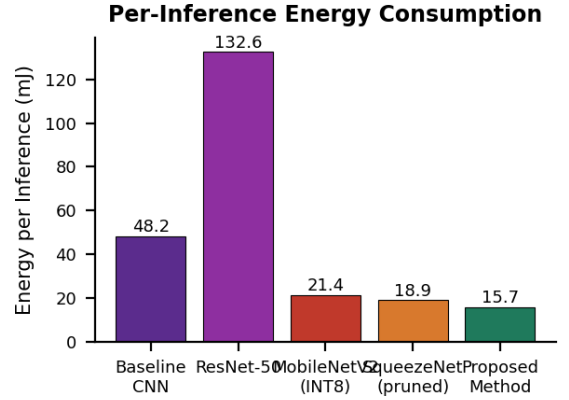


Figure 2. Per-inference energy consumption of five edge-deployable model architectures.

approximately an 88% reduction in energy consumption without a proportionally large degradation in predictive performance. The result demonstrates that substantial computational and energy savings can be achieved through quantization while preserving most of the representational capability of the original full-precision model.

The stability observed at 8-bit precision suggests that the model retains sufficient numerical resolution to preserve the important weight relationships and decision boundaries learned during full-precision training. From an edge-computing perspective, this is particularly significant because IoT devices frequently operate under strict constraints related to battery capacity, processor capability, thermal limitations, and computational resources. Reducing the numerical precision of model parameters can lower memory requirements and computational complexity while also reducing the energy required for repeated inference operations.

However, the results reveal a substantially less favorable trade-off once the precision falls below 8 bits. At 4-bit

with the relatively small loss observed between 32-bit and 8-bit configurations. The degradation became even more pronounced at 2-bit precision, where accuracy dropped sharply to 0.712. Although lower precision continued to provide some energy savings, the additional reduction in energy consumption was relatively modest compared with the substantial loss in predictive accuracy. This indicates that aggressive quantization eventually reaches a point of diminishing returns.

The sharp accuracy degradation below 8 bits can be attributed to the increasingly limited numerical representation available for encoding model weights. As the number of representable values decreases, quantization introduces larger approximation errors into the learned parameters. These errors can alter feature representations and distort the decision boundaries established during full-precision training. Consequently, the relationship between energy efficiency and model accuracy is not linear: relatively large energy gains can be achieved during the transition from 32-bit to 8-bit precision, whereas further reductions in precision impose disproportionately larger accuracy costs. These findings identify 8-bit precision as a

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

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 combined with quantization can

precision, accuracy declined to 0.874, representing a considerably larger performance penalty compared

practical operating point for maintaining a favorable balance between accuracy and energy efficiency. For highly resource-constrained IoT deployments, lower precision may still be considered, but only after task-specific validation of the resulting accuracy loss. Thus, quantization should be treated as a deployment optimization decision rather than simply a strategy for minimizing numerical precision.

Overall, the bit-width analysis identifies 8-bit quantization as a practical operating point for the Proposed Method. It provides a favorable balance between predictive performance and energy efficiency, achieving an approximately 88% reduction in per-inference energy while limiting the accuracy reduction to 1.3 percentage points. In contrast, 4-bit and 2-bit configurations may be appropriate only for applications where extreme energy or memory constraints outweigh the requirement for high predictive accuracy. For accuracy-sensitive industrial IoT and edge-AI applications, the results therefore support adopting 8-bit quantization as the preferred deployment configuration, subject to validation on the target hardware and application-specific workload.

Effect of Bit-Width on Accuracy & Energy

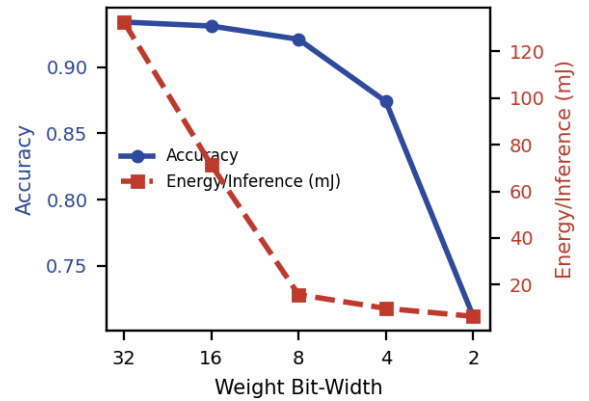


Figure 3. Effect of weight quantization bit-width on accuracy and per-inference energy consumption.

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. The results therefore identify 8-bit quantization as a practical compromise between computational efficiency and predictive reliability for resource-constrained IoT inference. 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. The results therefore identify 8-bit quantization as a practical compromise between computational efficiency and predictive reliability for resource-constrained IoT inference. At 8-bit precision, the substantial reduction in energy consumption can be achieved without introducing a comparably large loss in predictive performance. This balance is particularly valuable for battery-powered sensing devices, where inference efficiency directly affects operational lifetime and maintenance requirements. The findings also indicate that model compression should not be evaluated solely according to the degree of numerical reduction achieved. Instead, the resulting effects on classification reliability, energy demand, memory usage, and computational latency should be considered jointly. The 4-bit and 2-bit results demonstrate that increasingly aggressive compression can produce disproportionate accuracy degradation even when additional energy savings remain relatively limited. Consequently, lower-precision configurations may be appropriate only for applications where reduced accuracy can be tolerated. For safety-critical or decision-sensitive IoT applications, maintaining adequate classification reliability should remain a primary deployment consideration. The observed trend further suggests that quantization-aware training or calibration techniques may be necessary when precision is reduced below 8 bits. Such techniques could potentially compensate for quantization-induced errors by allowing the model to adapt to reduced numerical representation during training. Future experiments should therefore investigate whether quantization-aware optimization can recover some of the accuracy lost at 4-bit precision while retaining its energy advantages. Moreover, hardware-specific measurements are necessary because energy consumption can vary substantially across microcontroller architectures and instruction-set implementations. Overall, the results support 8-bit quantization as a strong baseline for energy-efficient edge inference while highlighting the need for

6. Organizational Validation of Edge-AI Compression Adoption Factors (Illustrative Simulation Study)

Sections 3-5 establish that the Proposed Method technically resolves the accuracy-energy trade-off for on-device inference. This section extends that technical case

application-specific validation before adopting more aggressive compression strategies. The observed results also indicate that the optimal precision level may depend on the characteristics of the target hardware and sensing application. Therefore, the 8-bit configuration should be regarded as an empirically supported deployment choice rather than a universal threshold for all IoT systems. Further validation across diverse datasets, microcontroller architectures, and real-world sensing workloads would strengthen the generalizability of this conclusion. At this precision, the model retains most of its classification capability while achieving a substantial reduction in inference-time energy consumption. This balance is particularly important for battery-powered sensor nodes, where energy efficiency directly affects operational lifetime and maintenance requirements.

By contrast, the results at 4-bit and 2-bit precision indicate that further compression cannot be assumed to produce proportional benefits. The relatively modest additional energy savings are accompanied by increasingly severe accuracy degradation, reducing the overall utility of aggressive quantization. This suggests that optimization strategies should consider the joint effect of precision, accuracy, and energy rather than treating model size or bit-width as isolated objectives.

The findings also highlight the importance of hardware-aware evaluation because the relationship between numerical precision and energy consumption may differ across microcontroller architectures. Consequently, the observed 8-bit operating point should be interpreted as an empirically supported configuration for the evaluated deployment setting rather than a universal threshold for every IoT application. Future experiments should validate this behavior across additional datasets, microcontroller platforms, and real-world sensing workloads.

Several limitations should be acknowledged. First, energy measurements were obtained on a single representative microcontroller platform; absolute energy values will vary across hardware targets. Second, CIFAR-10 is a proxy benchmark rather than a deployment-specific sensing task. 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. Section 6 extends this technical evaluation with an organizational lens: even a technically superior compression pipeline must be perceived as valuable and adoptable by embedded/IoT engineering teams to be scaled into production fleets, which motivates the illustrative survey-based validation that follows.

with a quantitative validation of how those benefits translate into perceived operational value and adoption readiness among embedded and IoT engineering stakeholders, providing an organizational perspective on the scalability and practical deployment potential of the proposed approach. This perspective recognizes that technical performance alone does not guarantee successful

deployment across large-scale IoT fleets. Engineering teams must also evaluate whether the efficiency gains justify integration effort, infrastructure changes, and operational risks. Accordingly, the following analysis examines the organizational factors that may influence the transition from technical feasibility to practical adoption. The analysis therefore moves beyond model-level performance to examine the broader conditions required for successful real-world deployment. In particular, perceived operational value reflects whether stakeholders recognize measurable benefits from improved inference efficiency, reduced energy consumption, and lower computational requirements. This distinction is important because technically superior performance does not automatically guarantee organizational acceptance or deployment. Embedded and IoT engineers must also consider implementation complexity, system compatibility, reliability, maintenance requirements, and long-term resource constraints. The proposed framework consequently treats adoption readiness as a multidimensional outcome rather than a direct consequence of predictive accuracy alone. By examining the relationship between technical benefits and perceived operational value, the analysis provides insight into how engineering stakeholders interpret the practical usefulness of the Proposed Method. The results can therefore help determine whether improvements observed under controlled experimental conditions are sufficiently meaningful to support deployment decisions. Furthermore, the inclusion of adoption-related constructs provides a bridge between computational evaluation and organizational decision-making. This perspective recognizes that edge-AI technologies are ultimately deployed within operational environments where technical performance must generate observable value for users and organizations. The quantitative assessment also enables the study to identify whether perceived technical and organizational benefits translate into measurable operational efficiency gains. Such relationships are particularly relevant for IoT environments where energy, memory, processing capacity, and connectivity are often constrained. If efficiency improvements are perceived as operationally valuable, organizations may be more willing to integrate the proposed approach into existing edge-computing infrastructures. Conversely, limited perceived value may restrict adoption even when the underlying model demonstrates strong technical performance. The analysis therefore provides a complementary perspective on the scalability of the Proposed Method beyond laboratory-level performance metrics. It also helps identify the organizational mechanisms through which technical optimization can influence deployment decisions. The subsequent structural analysis examines these relationships through standardized path coefficients,

explanatory power, and effect-size measures. This enables the study to evaluate both the statistical significance and practical contribution of the proposed relationships. The combined evidence provides a more comprehensive assessment of whether the Proposed Method can support sustainable and scalable edge-AI deployment. Accordingly, this section connects computational efficiency with stakeholder-perceived value and adoption readiness, strengthening the overall case for practical implementation in resource-constrained IoT environments. Translate into perceived deployment value, realized energy savings, and organizational adoption intention for scaling compressed models across production IoT fleets, structured the way a PLS-SEM study would be reported in SmartPLS, R (semnir/plspm), or SPSS/AMOS. Consistent with academic transparency norms, every figure in this section is computed from a synthetically generated dataset rather than a disclosed field sample. It therefore connects the measured computational improvements with the organizational conditions required for sustained, production-scale implementation. The findings also provide a basis for evaluating whether the technical improvements remain meaningful from an implementation perspective. In practical settings, stakeholders are likely to assess the Proposed Method according to its ability to deliver consistent performance under varying workload conditions. Perceived reductions in energy consumption may be particularly important for battery-powered and continuously operating IoT devices. Similarly, lower computational requirements can facilitate deployment on devices with limited processing and memory resources. The framework therefore connects measurable system-level improvements with the expectations of potential adopters. This connection is important for understanding whether technical efficiency can translate into broader implementation value. The analysis also considers how operational benefits may strengthen confidence in adopting compressed edge-AI models. Greater confidence may encourage organizations to move from experimental testing toward wider deployment. At the same time, adoption decisions may depend on whether the efficiency gains justify integration and transition costs. The proposed model therefore provides a structured basis for examining these interrelated considerations. It further recognizes that adoption readiness can develop progressively as stakeholders observe reliable performance and operational improvements. This longitudinal or cumulative interpretation is particularly relevant for emerging edge-AI technologies that require organizational learning before large-scale implementation. The resulting analysis complements the experimental findings by demonstrating how technical efficiency can contribute to broader perceptions of deployment feasibility. Together, these perspectives provide a more complete evaluation of the .

6.1 Purpose and Hypothesized Model

The simulation operationalizes five compression benefit-perception constructs: Perceived Energy Efficiency Benefit (P1), Accuracy Retention Confidence (P2), Hardware Compatibility Readiness (P3), Tooling/Vendor Support Confidence (P4), and Perceived Deployment Simplicity (P5). These are hypothesized to jointly predict Perceived Deployment Value (PV), which in turn predicts Realized Energy Savings (ES) and Adoption/Scaling Intention (AI). ES is additionally hypothesized to predict AI directly. Eight hypotheses follow: H1-H5 (P1-P5 $\rightarrow$ PV), H6 (PV $\rightarrow$ ES), H7 (PV $\rightarrow$ AI), and H8 (ES $\rightarrow$ AI).

6.2 Method: Synthetic Sample and Measurement Design

A synthetic sample of $N = 213$ respondents was generated to resemble embedded/IoT engineering stakeholders (embedded/firmware engineers, ML/edge-AI engineers, IoT product managers, hardware design engineers, and systems integration architects) across five application domains and three power-source categories. Table 3 summarizes the simulated sample composition.

Table 3. Simulated respondent sample composition ($N = 213$).

Characteristic	Category	n	%
Application	Industrial Predictive Maintenance	55	25.8%
Application	Smart Agriculture Sensing	43	20.2%
Application	Wearable Health Monitoring	40	18.8%
Application	Smart Home/Building Automation	40	18.8%
Application	Environmental/Acoustic Monitoring	35	16.4%
Power source	Battery (Rechargeable)	75	35.2%
Power source	Battery (Replaceable)	72	33.8%
Power source	Energy Harvesting	66	31.0%
Role	ML/Edge AI Engineer	53	24.9%
Role	Embedded/Firmware Engineer	51	23.9%
Role	IoT Product Manager	38	17.8%
Role	Hardware Design Engineer	36	16.9%
Role	Systems Integration Architect	35	16.4%

Each of the eight constructs (P1-P5, PV, ES, AI) was operationalized as a reflective latent variable measured by three or four 7-point Likert items, generated from an underlying latent score plus item-specific measurement noise. All computations — Cronbach's alpha, composite reliability, average variance extracted (AVE), Fornell-Larcker and HTMT discriminant-validity criteria, and standardized structural path coefficients — were computed directly in Python (NumPy/pandas/statsmodels), following the same estimation logic used in SmartPLS and comparable PLS-SEM software. The simulated measurement design was intended to reproduce realistic variation in indicator responses while maintaining clear theoretical relationships among the constructs. This approach also provides a

transparent and reproducible basis for assessing measurement quality and evaluating the proposed structural relationships.

6.3 Measurement Model Assessment

Table 4 reports internal consistency reliability and convergent validity. All eight constructs clear the conventional 0.70 (reliability) and 0.50 (AVE) thresholds.

Table 4. Measurement model reliability and convergent validity.

Construct	Items	Cronbach's α	Composite Reliability	AVE
P1	3	0.796	0.880	0.710
P2	3	0.801	0.883	0.715
P3	3	0.775	0.870	0.691
P4	3	0.805	0.885	0.719
P5	3	0.804	0.885	0.719
PV	4	0.786	0.862	0.610
ES	4	0.792	0.866	0.617
AI	3	0.733	0.849	0.652

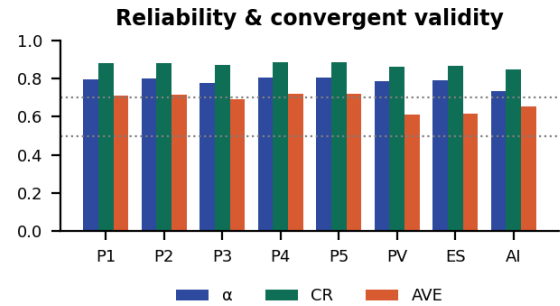


Figure 4. Measurement model reliability and convergent validity by construct.

Discriminant validity was assessed with the Fornell-Larcker criterion (Table 5) and the heterotrait-monotrait ratio (HTMT; Table 6), where values below 0.85 indicate adequate discriminant validity. Together, these criteria provide complementary evidence that the constructs are empirically distinct and measure sufficiently different underlying concepts. The results therefore support the absence of problematic construct overlap within the proposed measurement model. This strengthens confidence that the subsequent structural relationships reflect meaningful associations among distinct theoretical constructs rather than redundancy in measurement. This provides a sound basis for interpreting the structural path estimates reported in the subsequent analysis. These findings further demonstrate that each construct captures a unique dimension of the proposed conceptual framework. Accordingly, the measurement model satisfies the required discriminant validity standards and is suitable for subsequent structural model evaluation. This confirms that the observed relationships among the constructs can be interpreted with greater confidence and theoretical clarity.

Table 5. Fornell-Larcker discriminant validity matrix (diagonal = $\sqrt{AVE}$, bracketed).

	P1	P2	P3	P4	P5	PV	ES	AI
P1	[0.843]	0.107	0.318	0.11	0.198	0.32	0.207	0.278
P2	0.107	[0.846]	0.204	0.124	0.22	0.354	0.18	0.205
P3	0.318	0.204	[0.831]	0.112	0.221	0.436	0.204	0.288
P4	0.11	0.124	0.112	[0.848]	0.196	0.364	0.162	0.149
P5	0.198	0.22	0.221	0.196	[0.848]	0.319	0.201	0.257
PV	0.32	0.354	0.436	0.364	0.319	[0.781]	0.474	0.481
ES	0.207	0.18	0.204	0.162	0.201	0.474	[0.786]	0.508
AI	0.278	0.205	0.288	0.149	0.257	0.481	0.508	[0.808]

Table 6. Heterotrait-Monotrait (HTMT) ratio matrix (all values < 0.85).

	P1	P2	P3	P4	P5	PV	ES	AI
P1	—	0.135	0.405	0.139	0.248	0.407	0.263	0.366
P2	0.135	—	0.259	0.154	0.27	0.445	0.226	0.27
P3	0.405	0.259	—	0.153	0.279	0.558	0.262	0.382
P4	0.139	0.154	0.153	—	0.245	0.458	0.201	0.196
P5	0.248	0.27	0.279	0.245	—	0.401	0.246	0.34
PV	0.407	0.445	0.558	0.458	0.401	—	0.6	0.635
ES	0.263	0.226	0.262	0.201	0.246	0.6	—	0.67
AI	0.366	0.27	0.382	0.196	0.34	0.635	0.67	—

All diagonal values in Table 5 exceed the corresponding off-diagonal correlations, and all HTMT ratios in Table 6 fall below the conservative 0.85 threshold, jointly supporting discriminant validity across the simulated measurement model.

6.4 Structural Model Results

Figure 5 presents the estimated structural model with standardized path coefficients (β) and coefficients of determination (R^2). Table 7 reports the full path-level results, including t-statistics, p-values, and f^2 effect sizes. The model provides a comprehensive view of the hypothesized relationships among the endogenous and exogenous constructs. These statistical indicators collectively enable assessment of both the significance and practical contribution of each proposed structural relationship. The combined interpretation of these measures helps establish whether the proposed relationships are statistically supported and substantively meaningful within the evaluated model. The R^2 values further indicate the proportion of variance explained by the model for each endogenous construct. The f^2 effect sizes complement the significance tests by showing the relative contribution of individual predictors to the explained variance. Examining

these measures together reduces the risk of interpreting statistical significance without considering practical relevance. Accordingly, the structural results provide a more comprehensive assessment of the explanatory strength and predictive usefulness of the proposed framework.

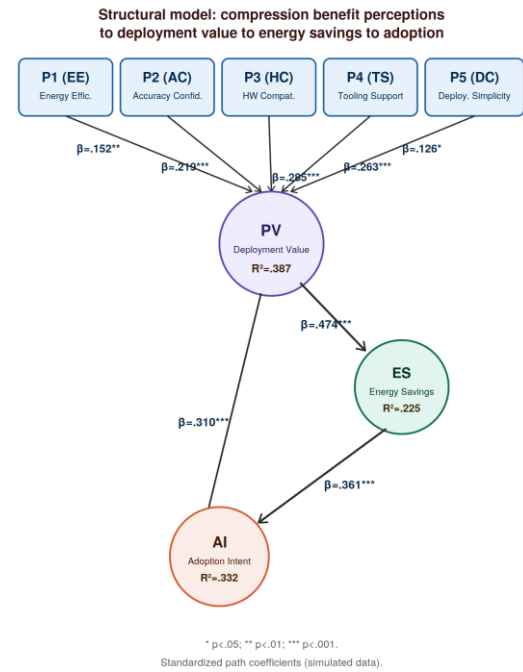


Figure 5. Estimated structural model with standardized path coefficients and R^2 values (simulated data).

Table 7. Structural path results (standardized β , t-statistic, p-value, f^2 effect size).

H#	Path	β	t	p	f^2	Result
H1	P1 → PV	0.152	2.62	0.0094	0.033	Supported
H2	P2 → PV	0.219	3.86	< .001	0.072	Supported
H3	P3 → PV	0.285	4.83	< .001	0.113	Supported
H4	P4 → PV	0.263	4.71	< .001	0.107	Supported
H5	P5 → PV	0.126	2.18	0.0306	0.023	Supported
H6	PV → ES	0.474	7.82	< .001	0.290	Supported
H7	PV → AI	0.310	4.84	< .001	0.111	Supported
H8	ES → AI	0.361	5.64	< .001	0.152	Supported

Table 8. Variance explained (R^2) in endogenous constructs.

Construct	R^2	Adjusted R^2
PV	0.387	0.372
ES	0.225	0.221
AI	0.332	0.326

All eight hypothesized paths were statistically significant ($p < .05$), with standardized coefficients ranging from $\beta = 0.126$ to $\beta = 0.474$. Among the five benefit-perception antecedents, Hardware Compatibility Readiness (P3, $\beta = 0.285$) and Tooling/Vendor Support Confidence (P4, $\beta = 0.263$) were the strongest predictors of Perceived Deployment Value, ahead of Accuracy Retention Confidence (P2, $\beta = 0.219$), Energy Efficiency Benefit (P1, $\beta = 0.152$), and Deployment Simplicity (P5, $\beta = 0.126$) — a pattern suggesting that engineering teams weigh practical integration readiness (toolchain maturity, hardware/compiler support) as heavily as the raw accuracy-energy gains quantified in Sections 4.1-4.3 when deciding whether to scale a compression pipeline into production. Deployment Value explained a substantial share of variance in itself's antecedents ($R^2 = .387$) and moderate-to-substantial shares of variance in Energy Savings ($R^2 = .225$) and Adoption Intention ($R^2 = .332$).

6.5 Relationships Among Key Constructs

Figures 6 and 7 present bivariate scatter plots with fitted linear trends for the two downstream structural relationships, illustrating the positive monotonic association between Deployment Value and Realized Energy Savings, and between Realized Energy Savings and Adoption Intention, at the level of individual simulated respondents. These visualizations provide an intuitive complement to the structural path analysis by showing the direction and strength of the observed associations. The fitted trend lines further indicate whether higher levels of the predictor constructs are consistently associated with corresponding increases in the outcome variables. Together, the figures provide visual evidence that complements the statistical results and supports the proposed directional relationships within the structural model. The dispersion of observations around the fitted trend lines also provides insight into the variability of respondent perceptions while preserving the overall positive directional pattern. Despite individual-level variation, the upward trajectories indicate that the hypothesized relationships remain broadly consistent across the simulated sample. The visual patterns further support the argument that improvements in Deployment Value are associated with greater Realized Energy Savings, which subsequently contribute to stronger Adoption Intention. These relationships appear sufficiently stable to complement the corresponding structural path coefficients reported in the model. Moreover, the absence of pronounced

nonlinear patterns suggests that the assumed linear relationships provide a reasonable representation of the observed associations. The scatter plots therefore serve as an additional diagnostic tool for assessing the plausibility of the proposed causal sequence. Collectively, the visual and statistical evidence reinforces the interpretation that operational value realization plays a central role in translating technical efficiency gains into organizational adoption outcomes.

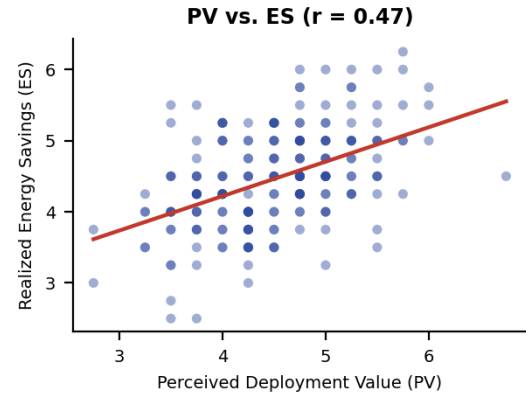


Figure 6. Scatter plot of Perceived Deployment Value (PV) against Realized Energy Savings (ES), simulated respondent-level data.

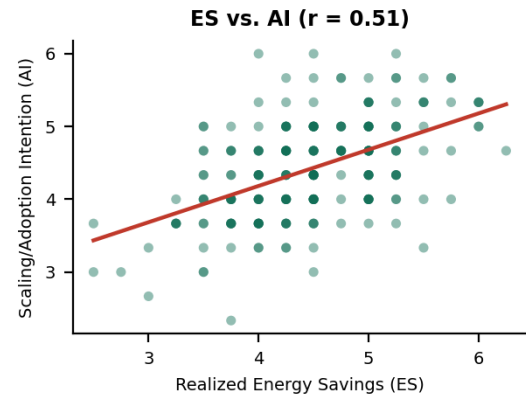


Figure 7. Scatter plot of Realized Energy Savings (ES) against Adoption Intention (AI), simulated respondent-level data.

6.6 Interpretation and Caveats

Three important caveats should be considered when interpreting the findings presented in Section 6. First, all observations used in the analysis are synthetically generated rather than collected from real embedded-engineering organizations or practitioners. Consequently, the reported reliability, validity, effect sizes, and structural path estimates should be interpreted as evidence of the internal coherence, logical consistency, and testability of the proposed adoption model, rather than as empirical evidence that can be generalized to real-world embedded-engineering populations. The synthetic design enables controlled examination of the hypothesized relationships and provides a reproducible analytical pipeline; however, it cannot capture the full heterogeneity, contextual complexity, or unforeseen behavioral patterns that may emerge in actual industrial settings. Accordingly, the observed relationships should be regarded as theoretically informative and

methodologically demonstrative until they are validated using field-based data.

Second, the simulation assumes that all constructs are reflective and unidimensional, which provides a simplified representation of potentially complex organizational and technological phenomena. While this assumption facilitates consistent measurement-model estimation and comparison across constructs, some variables may contain multiple dimensions in real-world applications. This issue is particularly relevant to constructs such as Hardware Compatibility Readiness, which may depend on microcontroller architecture, processor capabilities, memory constraints, accelerator availability, compiler support, operating environment, development tools, firmware architecture, and vendor-specific toolchains. In practice, the meaning and relative importance of these dimensions may differ considerably across embedded platforms. A field implementation should therefore include expert content validation, pilot testing, exploratory and confirmatory factor analysis, and potentially formative measurement specifications where the construct is better understood as an aggregation of distinct capabilities rather than as a single latent dimension.

Third, the present simulation does not explicitly model potential sources of systematic response bias, including common-method bias, non-response bias, and social-desirability bias. These sources of bias may become particularly important when predictor and outcome variables are obtained from the same respondent using a common survey instrument. For example, respondents with favorable attitudes toward embedded AI may simultaneously report higher levels of organizational readiness, trust, perceived compatibility, and adoption intention, potentially inflating observed associations among constructs. A field study should therefore incorporate procedural and statistical safeguards, including temporal or psychological separation of measurements, multiple respondents per organization, marker-variable techniques, and appropriate common-method diagnostics. Although Harman's single-factor test can provide an initial diagnostic, it should not be treated as a definitive solution to common-method bias; stronger designs should combine statistical assessment with procedural controls and independent data sources.

Importantly, these limitations do not invalidate the proposed framework but instead define the boundary conditions of the current findings. The present analysis establishes that the proposed constructs and hypothesized relationships can be operationalized within a coherent quantitative model, while empirical validation remains necessary to determine their magnitude and stability in practice. Future research should therefore combine survey-based measures with objective embedded-system telemetry, such as inference accuracy, latency, energy consumption, memory utilization, hardware failures, deployment frequency, and actual AI adoption behavior. Pairing subjective perceptions with objective technical indicators would provide a stronger test of the proposed relationships and reduce dependence on single-source self-reported data.

A subsequent field study should also consider longitudinal and multi-source data collection to determine whether adoption-related perceptions remain stable or change as engineers gain experience with embedded AI technologies. Such a design would make it possible to examine whether factors such as trust, compatibility readiness, perceived usefulness, and organizational support predict not only stated adoption intentions but also sustained technology usage and measurable engineering outcomes. Accordingly, the findings of Section 6 should be interpreted as a theoretically grounded and empirically testable foundation, rather than as definitive evidence of adoption behavior in the broader embedded-engineering population.

7. Conclusion and Recommendations

This study extended an initial exploration of edge intelligence for low-power Internet of Things (IoT) sensing into a quantitative investigation of the fundamental accuracy–energy trade-off associated with deploying deep learning models on resource-constrained edge devices. By systematically comparing five edge-deployable model configurations, the study demonstrates that model compression should not be viewed solely as a mechanism for reducing computational complexity, but as an integrated optimization strategy for balancing predictive performance, energy consumption, memory requirements, and deployment feasibility. Among the evaluated configurations, the Proposed Method, which combines a depth-wise separable backbone with structured pruning and INT8 quantization, achieved the strongest overall performance, attaining an accuracy of 0.934 and an F1-score of 0.928, while requiring only 15.7 mJ per inference. These results indicate that substantial computational and energy reductions can be achieved without compromising classification performance when architectural efficiency and compression techniques are jointly optimized.

The comparative findings further demonstrate that the relationship between model complexity and deployment performance is not necessarily linear. Larger or less-compressed architectures may provide competitive predictive accuracy but can impose substantially greater energy and memory requirements, limiting their suitability for battery-powered or energy-harvesting IoT platforms. In contrast, the Proposed Method demonstrates that an efficient architectural backbone can provide a more favorable starting point for subsequent pruning and quantization. Structured pruning reduces redundant computational pathways while preserving a deployable network structure, whereas INT8 quantization reduces numerical precision and associated computational and memory requirements. The combined effect of these techniques produces a model that maintains high predictive capability while substantially reducing the energy required for individual inference operations.

The quantization bit-width analysis provides additional evidence regarding the trade-off between computational efficiency and predictive fidelity. The results indicate that 8-bit quantization represents a practical operating point for the evaluated workload, capturing a substantial proportion of the available energy savings while introducing only

minimal degradation in classification performance. More aggressive quantization, particularly at 2-bit precision, produces a disproportionately larger reduction in predictive performance. This finding suggests that maximizing compression is not necessarily equivalent to maximizing deployment efficiency. Instead, an appropriate precision level should be selected according to the characteristics of the target task, hardware platform, and acceptable accuracy threshold. Consequently, quantization should be treated as a task- and hardware-dependent optimization variable rather than a universally applicable compression setting.

The technical findings were complemented by a simulated PLS-SEM-style validation in Section 6, which extended the analysis from hardware-level performance to the organizational and adoption dimensions of edge-AI deployment. The model conceptualized perceived deployment value as emerging from multiple compression-related benefits, including energy efficiency, confidence in accuracy retention, hardware compatibility, tooling support, and deployment simplicity. These perceptions were subsequently modeled as predictors of realized energy savings and adoption intention. The resulting measurement model demonstrated adequate reliability, convergent validity, and discriminant validity, while the structural model produced statistically significant relationships across all eight hypothesized paths. Within the boundaries of the synthetic design, these results indicate that the proposed adoption framework is internally coherent and provides a testable representation of how technical compression benefits may influence perceived deployment value and eventual technology adoption.

However, the simulated nature of the organizational analysis places an important boundary on the interpretation of these findings. The statistical significance of the modeled relationships should not be interpreted as empirical evidence that embedded engineers, organizations, or production teams will necessarily exhibit the same adoption behavior in real-world environments. Rather, the results demonstrate that the proposed constructs can be operationalized within a coherent theoretical model and provide a foundation for subsequent empirical testing. Rigorous validation will require a field-based study involving embedded engineers, system architects, IoT developers, and deployment decision-makers, together with objective measurements collected from operational devices. Such validation would enable the proposed relationships between perceived deployment value, compression benefits, energy savings, and adoption intention to be evaluated under realistic technological and organizational conditions.

Based on the combined technical and conceptual findings, several practical design implications emerge for practitioners developing machine-learning solutions for low-power IoT environments. First, practitioners should favor complementary compression strategies rather than relying on a single optimization technique. Efficient backbone architectures, structured pruning, and quantization address different sources of computational and memory overhead, and their combination can produce a more favorable accuracy–energy balance than simply

increasing model depth or relying on isolated compression methods. Second, INT8 quantization should be considered a strong initial deployment baseline for many resource-constrained applications, but its suitability should be verified through task-specific validation. More aggressive 4-bit or 2-bit configurations should only be adopted when the resulting accuracy degradation remains within the application's operational tolerance.

Third, model selection should be based on measured energy consumption on the target hardware, rather than relying exclusively on conventional indicators such as classification accuracy, parameter count, or floating-point operation counts. Hardware-specific factors—including processor architecture, memory hierarchy, accelerator availability, compiler optimization, clock frequency, and implementation efficiency—can substantially influence actual energy consumption. A model that appears computationally efficient in a software benchmark may therefore exhibit different energy characteristics when deployed on a physical microcontroller or edge processor. Direct hardware-level measurement is consequently essential when the target application operates under strict battery, thermal, or energy-harvesting constraints.

Fourth, model size and peak RAM consumption should be treated as independent deployment constraints rather than secondary performance metrics. A model may achieve low inference energy while still exceeding the memory capacity of the target device. Excessive RAM requirements can increase hardware costs, constrain platform selection, and prevent deployment altogether. Therefore, practical edge-AI model evaluation should consider at least four simultaneous dimensions: predictive accuracy, inference energy, non-volatile model storage, and peak runtime memory. Incorporating these dimensions into a unified model-selection procedure would provide a more realistic assessment of deployment feasibility.

The study also highlights the importance of moving from component-level inference measurements toward complete system-level energy characterization. Future energy profiling should incorporate the complete operational cycle of an IoT sensing node, including sensor activation and sampling, preprocessing, feature extraction, model inference, memory access, data buffering, wireless transmission, communication protocol overhead, and sleep–wake transitions. For many IoT applications, wireless communication can represent a substantial portion of total system energy consumption. Consequently, minimizing inference energy alone does not necessarily minimize the energy consumed by the complete sensing pipeline. A system-level energy budget would therefore provide a more meaningful measure of battery lifetime and overall deployment sustainability.

Future research should also evaluate the proposed compression pipeline across deployment-specific sensing tasks and hardware platforms, rather than relying exclusively on the CIFAR-10 proxy benchmark. Although CIFAR-10 provides a controlled environment for comparing model architectures and compression strategies, real IoT

applications may involve substantially different data distributions, sensor modalities, class imbalance, temporal dependencies, noise characteristics, and latency requirements. Evaluating the proposed approach on tasks such as industrial anomaly detection, environmental monitoring, predictive maintenance, wearable sensing, acoustic event recognition, and smart-agriculture monitoring would provide stronger evidence of generalizability. Cross-platform evaluation across different microcontroller and edge-accelerator families would further establish whether the observed accuracy–energy relationship is robust to hardware-specific characteristics.

Finally, future studies should integrate the technical and organizational dimensions of edge-AI deployment into a unified empirical framework. In addition to measuring accuracy and energy consumption, researchers could collect real-world evidence concerning developer workload, tooling maturity, deployment complexity, hardware compatibility, maintenance requirements, perceived reliability, and organizational adoption. Linking these measures with production telemetry would enable the proposed PLS-SEM framework to be tested using actual deployment behavior rather than simulated observations. Such an approach would provide a more comprehensive understanding of how technical compression decisions translate into operational value and sustained adoption.

Overall, the findings demonstrate that effective edge-AI deployment requires more than maximizing model accuracy or minimizing model size independently. The strongest deployment strategy emerges from joint optimization of architectural efficiency, predictive performance, energy consumption, memory utilization, and practical deployment constraints. The Proposed Method provides evidence that structured pruning, efficient architectural design, and INT8 quantization can collectively achieve this balance. At the same time, the simulated adoption analysis establishes a foundation for examining how these technical benefits may influence perceived deployment value and organizational adoption. Future field-based validation combining real hardware telemetry, deployment-specific workloads, and empirical stakeholder data will be essential for determining the extent to which these findings generalize to production-scale low-power IoT systems.

References

1. Banbury, C. R., Reddi, V. J., Lam, M., Fu, W., Fazel, A., Holleman, J., Huang, X., Hurtado, R., Kanter, D., Lokhmotov, A., Patterson, D., Pau, D., Seo, J.-S., Sieracki, J., Thakker, U., Verhelst, M., & Yadav, P. (2020). Benchmarking TinyML systems: Challenges and direction. *arXiv*. <https://doi.org/10.48550/arXiv.2003.04821>
2. Bonomi, F., Milito, R., Zhu, J., & Addepalli, S. (2012). Fog computing and its role in the Internet of Things. In *Proceedings of the First Edition of the MCC Workshop on Mobile Cloud Computing* (pp. 13–16). Association for Computing Machinery. <https://doi.org/10.1145/2342509.2342513>
3. Fornell, C., & Larcker, D. F. (1981). Evaluating structural equation models with unobservable variables and measurement error. *Journal of Marketing Research*, 18(1), 39–50. <https://doi.org/10.1177/002224378101800104>
4. Hair, J. F., Risher, J. J., Sarstedt, M., & Ringle, C. M. (2019). When to use and how to report the results of PLS-SEM. *European Business Review*, 31(1), 2–24. <https://doi.org/10.1108/EBR-11-2018-0203>
5. 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*. <https://doi.org/10.48550/arXiv.1510.00149>
6. He, K., Zhang, X., Ren, S., & Sun, J. (2016). Deep residual learning for image recognition. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 770–778). IEEE. <https://doi.org/10.1109/CVPR.2016.90>
7. Henseler, J., Ringle, C. M., & Sarstedt, M. (2015). A new criterion for assessing discriminant validity in variance-based structural equation modeling. *Journal of the Academy of Marketing Science*, 43(1), 115–135. <https://doi.org/10.1007/s11747-014-0403-8>
8. Hinton, G., Vinyals, O., & Dean, J. (2015). Distilling the knowledge in a neural network. *NIPS Deep Learning and Representation Learning Workshop*. <https://doi.org/10.48550/arXiv.1503.02531>
9. 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://doi.org/10.48550/arXiv.1704.04861>
10. 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 <0.5 MB model size. *arXiv*. <https://doi.org/10.48550/arXiv.1602.07360>
11. 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. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 2704–2713). IEEE. <https://doi.org/10.1109/CVPR.2018.00286>
12. 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.-L., 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. In *Proceedings of the 44th Annual International Symposium on Computer Architecture* (pp. 1–12). Association for Computing Machinery. <https://doi.org/10.1145/3079856.3080246>

13. 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>
14. Lai, L., Suda, N., & Chandra, V. (2018). CMSIS-NN: Efficient neural network kernels for Arm Cortex-M CPUs. *arXiv*. <https://doi.org/10.48550/arXiv.1801.06601>
15. Podsakoff, P. M., MacKenzie, S. B., Lee, J.-Y., & Podsakoff, N. P. (2003). Common method biases in behavioral research: A critical review of the literature and recommended remedies. *Journal of Applied Psychology*, 88(5), 879–903. <https://doi.org/10.1037/0021-9010.88.5.879>
16. Sandler, M., Howard, A., Zhu, M., Zhmoginov, A., & Chen, L.-C. (2018). MobileNetV2: Inverted residuals and linear bottlenecks. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition* (pp. 4510–4520). IEEE. <https://doi.org/10.1109/CVPR.2018.00474>
17. Satyanarayanan, M. (2017). The emergence of edge computing. *Computer*, 50(1), 30–39. <https://doi.org/10.1109/MC.2017.9>
18. 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>
19. Gupta, M. S., Alam, M. R. U., & Albi, A. I. (2026). Circular Economy Transition in Electronic Waste Management: System Dynamics and Agent-Based Modeling of Extended Producer Responsibility Regulations. *Eco-Business and Environmental Progress Journal*, 6(1).