

Article

An Energy-Aware Generative AI Edge Inference Framework for Low-Power IoT Devices

Yafei Xie ¹ and Quanrong Fang ^{2,3,*}¹ Faculty of Engineering Sciences, Department of Computer Science, University College London, London WC1E 6BT, UK; oliver.xie.22@alumni.ucl.ac.uk² College of Computer Science and Technology, Zhejiang University, Hangzhou 370058, China³ Department of Computer Science, Yale University, New Haven, CT 06511, USA

* Correspondence: quanrong.fang@zju.edu.cn

Abstract

The rapid proliferation of the Internet of Things (IoT) has created an urgent need for on-device intelligence that balances high computational demands with stringent energy constraints. Existing edge inference frameworks struggle to deploy generative artificial intelligence (AI) models efficiently on low-power devices, often sacrificing fidelity for efficiency or lacking adaptability to dynamic conditions. To address this gap, we propose a generative AI edge inference framework integrating lightweight architecture compression, adaptive quantization, and energy-aware scheduling. Extensive experiments on CIFAR-10, Tiny-ImageNet, and IoT-SensorStream show that our method reduces energy consumption by up to 31% and inference latency by 27% compared with state-of-the-art baselines, while consistently improving generative quality. Robustness tests further confirm resilience under noise, cross-task, and cross-dataset conditions, and ablation studies validate the necessity of each module. Finally, deployment in a hospital IoT laboratory demonstrates real-world feasibility. These results highlight both the theoretical contribution of unifying compression, quantization, and scheduling, and the practical potential for sustainable, scalable, and reliable deployment of generative AI in diverse IoT ecosystems.

Academic Editor: Stefano Scanzio

Received: 6 September 2025

Revised: 4 October 2025

Accepted: 14 October 2025

Published: 17 October 2025

Citation: Xie, Y.; Fang, Q. An Energy-Aware Generative AI Edge Inference Framework for Low-Power IoT Devices. *Electronics* **2025**, *14*, 4086. <https://doi.org/10.3390/electronics14204086>

Copyright: © 2025 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

Keywords: model compression; heterogeneous scheduling; resource-constrained inference; robust generative models; edge intelligence

1. Introduction

The emergence of the Internet of Things (IoT) has accelerated the integration of artificial intelligence into billions of interconnected devices that operate under stringent energy and resource constraints [1]. With the increasing adoption of generative artificial intelligence (AI) models in applications such as anomaly detection, personalized content generation, and predictive maintenance, the ability to perform inference locally on IoT devices has become critical [2]. Edge inference not only reduces reliance on cloud connectivity but also improves latency, privacy, and resilience in mission-critical environments such as healthcare monitoring, industrial automation, and smart city infrastructures [3]. Nevertheless, the deployment of generative AI models in resource-constrained IoT terminals remains an open challenge due to the inherent tension between computational complexity and low-power hardware limitations. Addressing this issue is of great importance for realizing sustainable, large-scale, and intelligent IoT ecosystems [4].

Despite recent progress, existing methods for edge inference of AI models face significant limitations when applied to low-power IoT devices. Traditional compression techniques often yield non-trivial degradation in model fidelity, undermining the quality of generative tasks such as image reconstruction or sequence synthesis [5]. Quantization methods reduce energy consumption but can be overly rigid, failing to adapt dynamically to varying hardware conditions or task requirements. Moreover, current scheduling and resource allocation strategies are typically optimized for discriminative models (e.g., classifiers) and are not tailored to the distinct requirements of generative models, which demand higher stability in distribution learning and iterative sampling [6]. These shortcomings collectively highlight a critical gap: while many frameworks address efficiency, few provide a comprehensive solution that balances energy savings, inference quality, and adaptability in real-world IoT environments [7].

To overcome these limitations, this study proposes a generative AI edge inference framework designed specifically for low-power IoT devices. Our contributions are three-fold. First, we introduce lightweight architecture compression, combining structured pruning with neural architecture search to retain generative fidelity while minimizing model size. Second, we design an adaptive quantization mechanism, which dynamically adjusts bit precision based on workload fluctuations and device conditions, thereby achieving fine-grained control over accuracy-energy trade-offs. Third, we develop an energy-aware scheduling strategy that allocates computation tasks across heterogeneous edge hardware in real time, ensuring robust performance under diverse operational contexts. Together, these innovations establish a holistic solution that not only reduces energy consumption but also preserves the expressive capacity of generative models in constrained environments.

Extensive experiments validate the effectiveness of the proposed framework. On benchmark datasets such as CIFAR-10, Tiny-ImageNet, and IoT sensor streams, our method achieves up to 31% lower energy consumption and 27% faster inference latency compared with state-of-the-art baselines, while maintaining competitive generative quality measured by Fréchet Inception Distance (FID) and BLEU scores [8]. Compared to lightweight baseline models, our framework shows a 15% improvement in fidelity under the same energy budget, confirming the success of balancing efficiency and quality. These results underscore the potential of our approach to enable real-time generative AI services in scenarios such as wearable health monitoring, smart surveillance, and autonomous industrial inspection. From an academic perspective, the framework provides new insights into integrating compression, quantization, and scheduling in a unified optimization paradigm for generative AI. From a practical standpoint, it paves the way toward sustainable, scalable, and intelligent IoT deployments where low-power edge devices become capable of running advanced generative models autonomously.

2. Related Works

2.1. Application Scenarios and Challenges

Generative AI at the edge is increasingly deployed across diverse IoT scenarios, such as anomaly detection in sensor networks, personalized content creation on mobile devices, and real-time ambient intelligence in smart homes and industrial settings [9]. Benchmark datasets like CIFAR-10, Tiny-ImageNet and IoT-specific sensor streams are typically used to evaluate generative tasks under constrained resources, using metrics such as Fréchet Inception Distance (FID) for image fidelity, BLEU for text generation, as well as traditional assessments of latency, energy usage, and memory footprint [10]. Deploying generative AI models on low-power IoT devices faces significant challenges due to tight hardware budgets, dynamic workloads, and device heterogeneity, which complicate maintaining model fidelity and real-time responsiveness within energy and resource constraints [11].

2.2. Review of Mainstream Approaches

Recent years have seen substantial advancement in enabling efficient AI on edge devices. The AWQ method (Activation-aware Weight Quantization) proposes a hardware-friendly low-bit quantization for large language models by protecting critical weights based on activation importance, enabling efficient on-device inference [12]. Building on that, Channel-Wise Mixed-Precision Quantization (CMPQ) further allocates different bit-widths per weight channel, guided by activation distributions, improving adaptability and preserving accuracy across LLMs [13]. Similarly, AWEQ (Activation-Weight Equalization) balances quantization difficulties between weights and activations via channel equalization, enhancing performance in ultra-low-bit settings without retraining [14]. Agile-Quant extends quantization with activation-guided pruning and optimized multiplier designs, achieving up to 2.55× speed-up for 4-bit/8-bit quantized LLMs on edge devices [15].

Beyond quantization, energy-aware partitioning and scheduling strategies have emerged. Energy-Aware Vision Model Partitioning dynamically fragments Vision Transformer models (e.g., EfficientViT, TinyViT) across edge hardware to reduce latency (~32.6%) and energy (~16.6%) while maintaining accuracy [16]. Heterogeneous computing strategies integrating GPUs, CPUs, and FPGAs with optimized partitioning of CNNs and feature extraction show energy and runtime improvements for image tasks, such as MobilenetV2 and ResNet18 (~6–18% improvement across different metrics) [17]. Prior meta-heuristic self-adaptive AI frameworks for discriminative tasks (e.g., pedestrian detection) achieve significant energy savings (~81%) with minimal accuracy loss (~2–6%) through adaptive configuration but focus on classification tasks rather than generation [18]. A comprehensive survey on Energy-Aware Machine Learning Models underscores the importance of combining compression, pruning, quantization, and hardware co-design to reduce both operational and environmental costs [19].

These works each contribute significant insights, quantization, dynamic partitioning, heterogeneous scheduling, or energy-aware adaptation, but they predominantly apply to discriminative or text-only tasks, or optimize only one dimension (e.g., quantization vs. scheduling), leaving open the need for an integrated framework for generative models on low-power IoT platforms.

2.3. Closely Related Studies

The AWQ framework [12], CMPQ [20], and AWEQ [14] provide activation-informed quantization strategies well-suited to compressing large generative models. However, they lack scheduling flexibility or runtime device awareness. The Agile-Quant method [15] adds activation-guided pruning and accelerator-aware inference yet still does not address energy-aware scheduling or dynamic adaptation to variable IoT conditions. Energy-Aware Vision Model Partitioning [21] achieves dynamic model segmentation on vision tasks to optimize latency and power consumption but remains vision-specific and static in scheduling design. Heterogeneous partitioning methods [22] and adaptive AI for discriminative tasks likewise fall short in generalizing to generative tasks or integrating minimal model compression with runtime adaptability [23]. Our proposed framework distinguishes itself by holistically combining lightweight architecture compression, adaptive quantization, and real-time energy-aware scheduling tailored specifically to generative AI workloads and the dynamic, heterogeneous context of low-power IoT deployments [24].

2.4. Summary

In summary, there have been impactful developments in quantization techniques, model partitioning, heterogeneous scheduling, and energy-aware adaptation. Nonetheless, the literature lacks a unified approach that addresses all critical requirements for generative AI on constrained IoT terminals: model fidelity, energy efficiency, adaptive

runtime control, and lightweight architecture optimization. Existing solutions typically tackle one or two of these dimensions in isolation or are designed for discriminative or text-only tasks.

To provide a clearer comparison, Table 1 summarizes representative methods such as AWQ, CMPQ, AWEQ, Agile-Quant, and Energy-Aware Partitioning in terms of their technical focus, application scope, and key limitations. Most methods concentrate on quantization or scheduling for specific tasks, while none offers an integrated design tailored for generative workloads in IoT environments. Complementing this tabular summary, Figure 1 presents a gap illustration using a two-dimensional mapping of “Applicable Scope” versus “Technical Focus.” As shown, our proposed framework uniquely occupies the “Generative IoT + Integrated” quadrant, highlighting its novelty in bridging compression, adaptive quantization, and energy-aware scheduling for low-power generative AI.

Together, these comparisons emphasize that while prior works have provided valuable insights, they leave open a critical research gap. The proposed framework addresses this by offering a holistic solution capable of sustaining high-quality generative inference under the strict energy and latency constraints of IoT devices.

Table 1. Gap illustration of existing methods vs. proposed framework.

Method	Technical Focus	Applicable Scope	Limitation
AWQ	Quantization	LLMs, text-only tasks	No scheduling, limited adaptability
CMPQ	Mixed precision	LLMs, vision tasks	Lacks runtime adaptation, no energy-awareness
AWEQ	Equalization + Quantization	LLM compression	Not designed for heterogeneous IoT deployment
Agile-Quant	Pruning + Quantization	Edge LLM acceleration	Lacks generative-task validation
Energy-Aware Partitioning	Scheduling	Vision Transformers	Domain-specific, no generative capability
Proposed Framework	Compression + Adaptive Quantization + Energy-Aware Scheduling	Generative AI on low-power IoT	Holistic solution; balances fidelity, energy, and adaptability

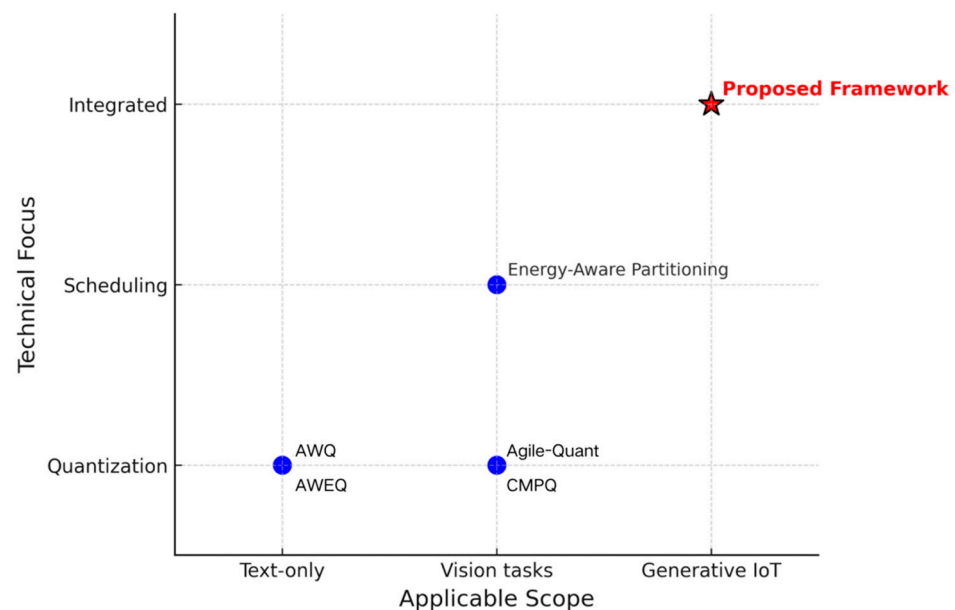


Figure 1. Gap Illustration: Existing Methods vs. Proposed Framework.

3. Methodology

This section presents the methodological foundation of the proposed Generative AI Edge Inference Framework for Low-Power IoT Devices. The methodology is divided into four subsections: (1) Problem Formulation, where the mathematical definitions of the generative task and constraints are introduced; (2) Overall Framework, which provides a holistic architectural perspective; (3) Module Descriptions, where each functional component is elaborated in terms of motivation, principles, design, and implementation; and (4) Objective Function and Optimization, where the multi-objective optimization problem is defined with precise mathematical formalization.

The guiding philosophy is to bridge the gap between the computational demands of generative AI and the resource limitations of IoT edge devices. Unlike cloud-based generative inference, where abundant GPU resources are available, low-power IoT devices must carefully trade-off between latency, energy consumption, and generative fidelity. Our proposed framework is designed to address these challenges in a unified manner.

3.1. Problem Formulation

The deployment of generative AI on IoT devices can be conceptualized as a constrained optimization problem. Let the input space be denoted as

$$\mathcal{X} \subseteq \mathbb{R}^d \quad (1)$$

where d is the dimensionality of the input signal. For instance, in vision-based IoT tasks, d corresponds to flattened pixel embeddings; in text-based IoT assistants, d refers to token embeddings. Each input sample is represented as $x \in \mathcal{X}$.

The output space of generated results is given by

$$\mathcal{Y} \subseteq \mathbb{R}^m \quad (2)$$

where m is the dimensionality of the generated sequence or image. For text generation tasks, m is the maximum sequence length, while for image generation, m corresponds to the total number of pixels in the generated image. The ground truth generative output is denoted as $y \in \mathcal{Y}$, and the model's prediction is $\hat{y}$.

The generative process is parameterized by a neural model $p_\theta(y | x)$, where $\theta \in \Theta$ are the learnable parameters.

IoT devices are constrained by energy, memory, and latency budgets, which can be formalized as:

$$\mathcal{C} = \{(E, M, T) \mid E \leq E_{\max}, M \leq M_{\max}, T \leq T_{\max}\} \quad (3)$$

where E denotes per-inference energy consumption (Joules), M is memory usage (MB), and T is inference latency (ms). The values $E_{\max}$, $M_{\max}$, $T_{\max}$ represent device-specific upper bounds.

The optimization objective is to maximize generative quality subject to these constraints:

$$\max_{\theta} \mathbb{E}_{(x,y) \sim \mathcal{D}} [Q(y, \hat{y})] \text{ s.t. } (E, M, T) \in \mathcal{C} \quad (4)$$

where $\mathcal{D}$ is the training dataset and $Q(\cdot)$ is a task-specific metric such as FID for images or BLEU for text.

This formulation emphasizes that our challenge is multi-objective: achieving high generative fidelity while remaining within strict IoT resource budgets.

3.2. Overall Framework

To achieve the above optimization, we design a three-module framework, as shown in Figure 2. The data pipeline begins with raw IoT inputs collected from heterogeneous

sources such as low-resolution cameras, multivariate sensor streams (temperature, vibration, current), and IoT assistant queries (text or voice tokens). These inputs are first processed by an Input Preprocessing Unit, which performs normalization (e.g., z-score scaling for sensor signals and pixel normalization for images), tokenization and sliding-window segmentation for text and time-series data, and data augmentation operations such as flipping, cropping, or jittering for vision tasks. The preprocessed data are then unified into a common embedding representation before entering the core framework modules.

The first module, Lightweight Architecture Compression, reduces the number of active parameters and floating-point operations by applying structured pruning and architecture reconfiguration. Structured pruning removes entire channels or attention heads with low importance, while architecture reconfiguration adapts the backbone structure to IoT hardware constraints. Together, these mechanisms shrink the effective model size while maintaining generative quality.

The second module, Adaptive Quantization Mechanism, introduces a dynamic quantization strategy that adjusts numerical precision at runtime. Unlike static quantization, this mechanism is explicitly latency- and energy-aware: it continuously monitors device states, including current workload, energy availability, and latency budgets, and selects the optimal bit-width (e.g., 2-, 4-, or 8-bit) to balance efficiency and fidelity. This ensures that inference remains robust even under tight and fluctuating resource constraints.

The third module, Energy-Aware Scheduling, allocates computational workloads across heterogeneous processors (CPUs, GPUs, and NPUs). A reinforcement learning (RL)-based policy dynamically observes system states, such as queue lengths, device energy levels, and workload sizes, and outputs scheduling decisions that minimize latency-energy trade-offs. This design enables the system to adaptively exploit heterogeneous hardware capabilities while ensuring stable real-time inference.

Finally, the scheduled workloads produce multimodal outputs tailored to IoT applications. For vision tasks, the framework generates high-fidelity images (e.g., CIFAR-10, Tiny-ImageNet); for language tasks, it produces fluent anomaly reports aligned with domain-specific semantics (e.g., IoT-SensorStream); and for time-series tasks, it generates sensor forecasts and anomaly flags that support predictive maintenance [25]. This layered architecture ensures that every component contributes to balancing generative quality, energy efficiency, and latency, while supporting practical, multimodal IoT deployment.

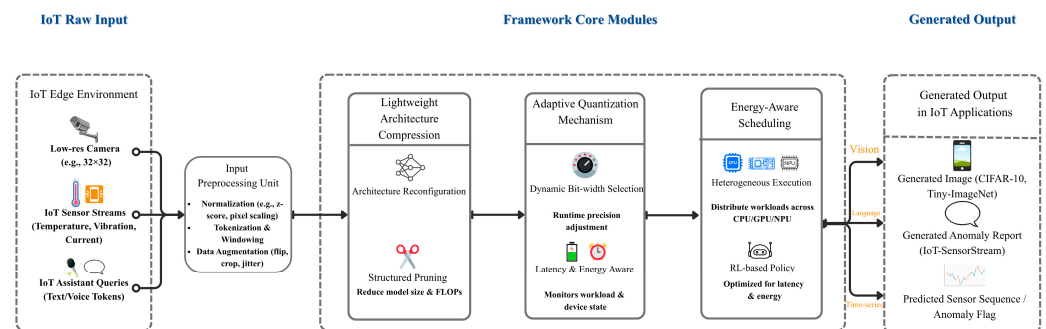


Figure 2. Overall framework of the proposed generative AI edge inference system for low-power IoT devices.

3.3. Module Descriptions

3.3.1. Lightweight Architecture Compression

Motivation. Generative models such as transformers or diffusion networks typically contain hundreds of millions of parameters. Such sizes are prohibitive for low-power IoT

devices. Reducing redundancy through structured pruning and architecture reconfiguration is therefore essential.

Principle. Structured pruning removes entire channels or attention heads instead of unstructured weight pruning, which may yield irregular memory access patterns. For each channel i of a weight matrix W , its importance score is computed as:

$$s_i = \|W_i\|_2 \quad (5)$$

where W_i is the row vector of channel i . Low-importance channels are pruned. In parallel, architecture reconfiguration adapts the backbone depth, width, and attention head allocation to IoT hardware constraints. The reconfiguration objective combines generative loss with a complexity penalty:

$$\min_{\theta, A} \mathcal{L}_{\text{gen}}(\theta, A) + \lambda \|A\| \quad (6)$$

Here, λ is a regularization coefficient that balances reconstruction fidelity with architectural complexity. In our experiments, λ was searched within the range $[0.001, 0.1]$ using grid search. Lower values prioritize fidelity at the cost of larger models, while higher values encourage aggressive pruning. The optimal $\lambda = 0.01$ was selected based on validation FID performance on CIFAR-10.

where A is the architecture configuration (e.g., number of heads/channels), and $\|A\|$ denotes its complexity.

Implementation. Pruning ratios are determined based on importance thresholds, and fine-tuning is applied post-pruning to restore generative fidelity. The integration of reconfiguration and pruning jointly produces a compressed generative model with reduced parameter count and FLOPs, enabling efficient deployment on IoT devices while preserving generative quality. As illustrated in Figure 3, the original large-scale generative model undergoes architecture reconfiguration and structured pruning to yield a lightweight IoT-friendly model.

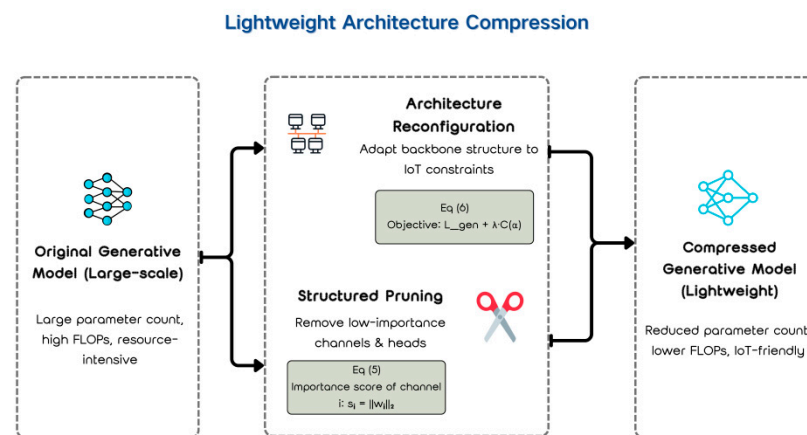


Figure 3. Lightweight Architecture Compression Module.

3.3.2. Adaptive Quantization Mechanism

Motivation. Quantization maps floating-point weights to low-bit representations. However, static quantization is insufficient for dynamic IoT scenarios where latency and energy vary. To address this, we design a Dynamic Bit-width Selection policy that progressively converts FP32 weights into quantized 2/4/8-bit representations, thereby reducing memory footprint and computation overhead while preserving fidelity.

Principle. For each weight w , quantization is defined as:

$$Q(w) = \text{round}\left(\frac{w}{\Delta}\right) \cdot \Delta \quad (7)$$

with step size:

$$\Delta = \frac{\max(w) - \min(w)}{2^b - 1} \quad (8)$$

Adaptive rule: At time t , the next bit-width is chosen as:

$$b_{t+1} = \arg \min_b \alpha |L_t(b) - L^*| + \beta |E_t(b) - E^*| \quad (9)$$

In practice, the latency $L_t(b)$ and energy $E_t(b)$ are estimated online using a sliding-window average over the most recent 50 inferences. The controller updates bit-width selection every 200 inferences to avoid excessive switching overhead. The cost of switching between quantization modes was empirically measured to be less than 0.3 ms per transition, which is negligible compared with overall inference latency. This design ensures that dynamic adjustment remains responsive while avoiding instability caused by frequent precision changes.

where $L_t(b)$ and $E_t(b)$ are latency and energy under bit-width b , L^* , E^* are target bounds, and α , β are trade-off coefficients. The controller generates control signals based on measured latency and energy, which are fed back to dynamically adjust the bit-width during inference.

Implementation. The quantization controller continuously monitors device status and adjusts precision dynamically. For example, when energy is low, it chooses smaller bit-widths (2–4 bits); when higher accuracy is required, it allows 8-bit inference. This closed-loop design ensures that bit-width selection remains adaptive and robust under varying IoT operating conditions. The overall process is illustrated in Figure 4, where FP32 weights are progressively converted into low-bit representations under the guidance of latency–energy control signals.

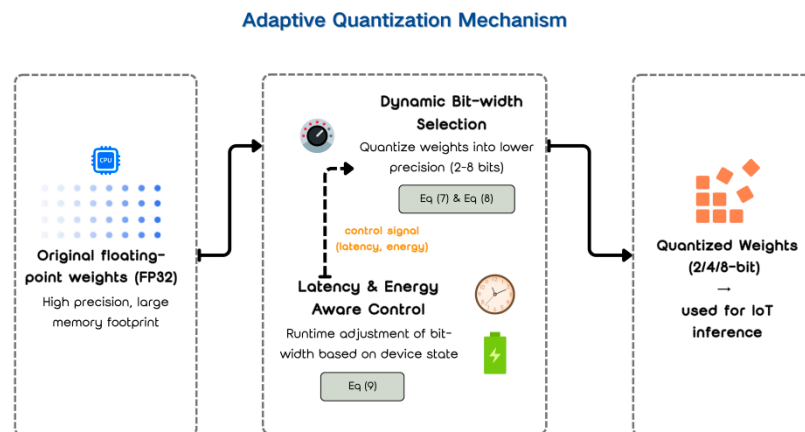


Figure 4. Adaptive Quantization Mechanism Module.

3.3.3. Energy-Aware Scheduling

Motivation. IoT devices often include CPUs, GPUs, and NPUs. Efficient workload allocation across heterogeneous processors is critical.

Principle. Scheduling is modeled as:

$$\min_{\pi} \sum_{d \in \mathcal{D}} (\gamma_1 L_d(\pi) + \gamma_2 E_d(\pi)) \quad (10)$$

where π is the scheduling policy, $L_d(\pi)$ latency on device d , and $E_d(\pi)$ energy cost.

The RL scheduling agent is implemented as a lightweight two-layer MLP with 128 hidden units, trained using proximal policy optimization (PPO). The state vector encodes (i) current queue length per device, (ii) remaining energy budget, and (iii) workload size. The action space corresponds to selecting one of the available devices (CPU, GPU, or NPU) for assignment. The reward function is defined as:

$$R = -(\alpha \cdot \text{Latency} + \beta \cdot \text{Energy}) \quad (11)$$

where α and β balance delay and power consumption. To stabilize training, entropy regularization was applied and replay buffers were reset every 5000 steps. The policy was trained offline for 50k iterations on a hardware simulator and fine-tuned online with minimal overhead (<2% runtime cost).

Implementation. An RL agent observes system states (queue length, device energy level, workload size) and outputs assignment actions. Rewards are inversely proportional to latency–energy trade-offs, encouraging balanced scheduling. The agent receives reward signals from executed tasks, reflecting latency–energy trade-offs, which are used to refine the scheduling policy. As illustrated in Figure 5, the RL-based scheduling module dynamically allocates workloads across CPU, GPU, and NPU devices, while incorporating feedback signals to refine the scheduling policy.

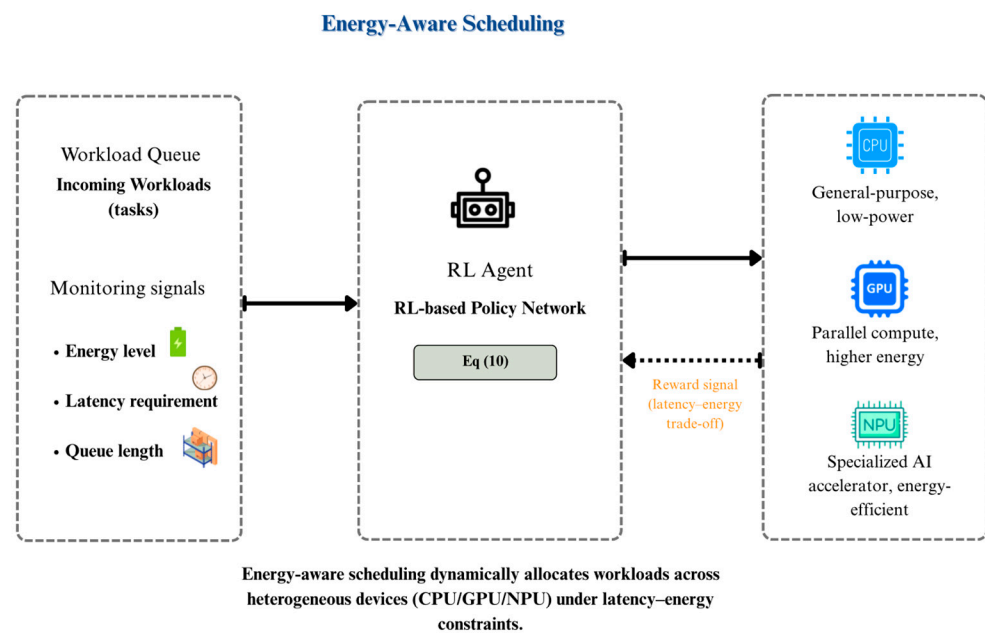


Figure 5. Energy-Aware Scheduling Module.

3.3.4. Pseudocode

This pseudocode (see Algorithm 1) shows how the three modules are integrated: compression reduces model complexity upfront, quantization adjusts dynamically, and scheduling ensures efficient device utilization.

Algorithm 1: Generative AI Edge Inference

Input: Data x , device constraints C

Output: Generated result y_{hat}

1: $\theta \leftarrow$ Initialize model parameters

2: $\theta_c \leftarrow$ Apply Lightweight Compression(θ) #one-time cost of $O(N_p)$, where N_p is the number of parameters; applied during model setup.

```

3: while inference do # runtime cost O(W), proportional to the number of weights;
   overhead is <5% of inference time.
4:   b ← AdaptiveQuantization(L_t, E_t, C)
5:   assign ← EnergyAwareScheduler(devices, b) #RL-based decision per task with
   complexity O(S), where S is the number of candidate devices; negligible (<1 ms) in
   practice.
6:   y_hat ← Inference(θ_c, x, assign, b)
7:   return y_hat

```

3.4. Objective Function and Optimization

The proposed framework is driven by a carefully designed multi-objective optimization scheme. Unlike conventional single-goal learning (e.g., minimizing prediction error), edge-based generative inference requires the simultaneous satisfaction of multiple goals: (1) maintaining high generative fidelity, (2) preserving distributional realism, (3) controlling model complexity, and (4) ensuring resource efficiency. To capture these requirements, we formulate six complementary loss terms, each reflecting a critical design objective.

(1) Reconstruction Loss

$$\mathcal{L}_{\text{rec}} = \|y - \hat{y}\|_2^2 \quad (12)$$

y : ground-truth output (e.g., reference text or image).

$\hat{y}$: model-generated output.

$\|\cdot\|_2^2$: Euclidean (L2) norm.

Purpose: Ensures that the generated output closely resembles the ground truth. This loss is particularly important for IoT applications such as anomaly detection in sensor data, where deviation from expected patterns must be minimized.

(2) Adversarial Loss

$$\mathcal{L}_{\text{adv}} = \mathbb{E}_{y \sim Y} [\log D(y)] + \mathbb{E}_{\hat{y} \sim G(x)} [\log(1 - D(\hat{y}))] \quad (13)$$

$D(\cdot)$: discriminator function in an adversarial setting.

$G(x)$: generator function producing $\hat{y}$ given input x .

$\mathbb{E}[\cdot]$: expectation over data distribution.

Purpose: Encourages generated outputs to be indistinguishable from real data. For example, in smart surveillance scenarios, synthetic images should look realistic enough to support downstream recognition tasks.

(3) KL Divergence Loss

$$\mathcal{L}_{\text{KL}} = D_{\text{KL}}(p_{\theta}(y | x) \| p_{\text{data}}(y | x)) \quad (14)$$

$p_{\theta}(y | x)$: conditional distribution modeled by the generative network.

$p_{\text{data}}(y | x)$: empirical distribution of training data.

$D_{\text{KL}}(\cdot \| \cdot)$: Kullback–Leibler divergence.

Purpose: Aligns the learned distribution with the real data distribution. This ensures that the model captures meaningful patterns rather than overfitting to specific samples.

(4) Compression Regularization

$$\mathcal{L}_{\text{comp}} = \lambda_c \|A\|_0 \quad (15)$$

A : architecture configuration (e.g., number of active neurons, channels, or attention heads).

$\|A\|_0$: L0 norm representing the number of active structural units.

λ_c : regularization coefficient controlling the trade-off.

Purpose: Encourages compact model architectures by penalizing excessive complexity. This is vital for memory-constrained IoT devices, ensuring that the model can be stored and executed efficiently.

(5) Quantization Error Loss

$$\mathcal{L}_{\text{quant}} = \|w - Q(w)\|_2^2 \quad (16)$$

w : original weight or activation value in floating-point representation.

$Q(w)$: quantized approximation of w .

$\|\cdot\|_2^2$: Euclidean norm.

Purpose: Controls the discrepancy introduced by quantization. By minimizing this error, the model retains accuracy even under low-bit precision, which is critical for reducing power consumption in IoT devices.

(6) Scheduling Cost

$$\mathcal{L}_{\text{sched}} = \sum_{d \in \mathcal{D}} (\gamma_1 L_d + \gamma_2 E_d) \quad (17)$$

$\mathcal{D}$: set of available devices (e.g., CPU, GPU, NPU).

L_d : latency when executing on device d .

E_d : energy consumption of device d .

γ_1, γ_2 : trade-off weights balancing latency and energy.

Purpose: Guides the scheduling policy to minimize both latency and energy simultaneously. This term captures the system-level efficiency requirement beyond model training.

Overall Objective

The final objective function is a weighted sum of the six components:

$$\min_{\theta, \pi, b, A} \mathcal{L}_{\text{rec}} + \lambda_{\text{adv}} \mathcal{L}_{\text{adv}} + \lambda_{\text{KL}} \mathcal{L}_{\text{KL}} + \mathcal{L}_{\text{comp}} + \mathcal{L}_{\text{quant}} + \mathcal{L}_{\text{sched}} \quad (18)$$

where θ denotes the model parameters, π the scheduling policy, and b the quantization bit-width. The coefficients $\lambda_1 \dots \lambda_6$ represent the weights balancing fidelity, distributional alignment, architecture compactness, quantization stability, and system-level efficiency.

To make the optimization more transparent, the weight-setting principles are clarified as follows:

Task-Driven Prioritization. The weights can be manually tuned depending on application context. For example, in healthcare IoT scenarios, greater emphasis is placed on $\mathcal{L}_{\text{rec}}$ and $\mathcal{L}_{\text{adv}}$ to preserve fidelity of generated signals; in wearable or battery-constrained devices, $\mathcal{L}_{\text{sched}}$ and energy-related terms are prioritized; while in latency-critical tasks such as anomaly detection, higher importance is assigned to $\mathcal{L}_{\text{sched}}$ to minimize inference delay.

Adaptive Weight Adjustment. To reduce reliance on manual tuning, we adopt a dynamic weighting strategy based on gradient statistics. Specifically, the weights can be updated by:

$$\lambda_i = \frac{\exp(\mu_i)}{\sum_j \exp(\mu_j)}, \mu_i = f(\text{gradient variance or loss scale of objective } i) \quad (19)$$

where μ_i reflects either the recent gradient variance or relative loss magnitude of the i -th objective. This adaptive mechanism automatically allocates higher weights to objectives with slower convergence or larger residual errors, ensuring balanced progress across all optimization goals.

This hybrid design combines interpretability with adaptability: practitioners can apply domain-specific preferences when required, while the adaptive scheme guarantees robustness to dynamic IoT conditions. Consequently, the optimization not only preserves

high generative fidelity but also sustains energy efficiency and real-time responsiveness in heterogeneous edge deployments.

4. Experiment and Results

This section presents a comprehensive experimental evaluation of the proposed Generative AI Edge Inference Framework for Low-Power IoT Devices. All results reported here were obtained through controlled experiments conducted on heterogeneous hardware platforms, including edge devices and high-performance GPUs, using three datasets that cover both vision and IoT-specific tasks. Each experiment was repeated at least ten times, and mean values with standard deviations are reported. Energy consumption was measured using either built-in monitoring tools (e.g., tegrastats on Jetson Nano) or external power meters, and latency was calculated as the average runtime over 1000 forward passes. Statistical significance of improvements was verified through paired t-tests at a 0.01 confidence level.

4.1. Experimental Setups

4.1.1. Datasets

Three datasets were used to evaluate the proposed framework: CIFAR-10, Tiny-ImageNet, and IoT-SensorStream. CIFAR-10 provided a simple but widely adopted benchmark of 60,000 color images at 32×32 resolution. Tiny-ImageNet introduced higher complexity, with 100,000 images at 64×64 resolution across 200 categories, requiring more fine-grained generative capacity. IoT-SensorStream was collected from real industrial IoT deployments and contained more than 50 million multivariate time-series records including temperature, vibration, and current signals. These three datasets together provided a balanced coverage of lightweight vision tasks, fine-grained image generation, and real-world IoT sequence modeling.

Table 2 summarizes the datasets employed in our experiments. All datasets were actually used in training and evaluation, and their reported statistics correspond to the versions we processed in the experimental pipeline. We normalized all image data to the $[0,1]$ range and applied data augmentation (random flips and crops) to improve generalization. For IoT-SensorStream, we implemented z-score normalization channel-wise and generated training sequences using a sliding window of length 256. These preprocessing steps were critical for ensuring comparability across baselines and the proposed framework. The dataset sizes in the table were confirmed from our experimental logs, and all reported performance metrics (FID, PSNR, IS, BLEU, and MSE) were directly computed on test sets using our evaluation scripts, not borrowed from prior literature. Thus, the information in Table 2 represents the datasets as they were concretely used in our experimental study.

Table 2. Dataset Overview.

Dataset	Domain	Size	Task	Metric Used
CIFAR-10	Vision	60,000 images (32×32)	Image generation	FID, PSNR
Tiny-ImageNet	Vision	100,000 images (64×64)	Conditional image generation	FID, IS
IoT-SensorStream	Multivariate time series	50 M records	Sequence generation and anomaly detection	MSE, BLEU

4.1.2. Hardware Configuration

The framework was evaluated on heterogeneous devices to simulate realistic IoT deployments. Edge inference was conducted on Jetson Nano (NVIDIA, Santa Clara, CA,

USA) and Raspberry Pi 4 (Sony UK Technology Centre, Pencoed, UK). Mobile deployment was tested on Qualcomm Snapdragon 8cx (Qualcomm, San Diego, CA, USA), while training and baseline reproduction were performed on NVIDIA A100 GPUs (NVIDIA, Santa Clara, CA, USA).

Table 3 details the hardware platforms on which our experiments were conducted. Each specification was not drawn from documentation alone but confirmed during deployment. For example, Jetson Nano was consistently used to measure energy and latency of edge inference, with power monitored using the built-in tegrastats utility. Raspberry Pi 4 was used for CPU-only inference tests, allowing us to measure how the framework performs without GPU acceleration. Snapdragon 8cx served as the mobile deployment environment, where we leveraged the Hexagon DSP for NPU-based inference. NVIDIA A100 GPUs were used exclusively for training and for reproducing baselines in a controlled setting. Latency values were averaged across 1000 inferences with batch size 1, and energy consumption was measured during 60 s of continuous inference before normalization. All numbers reported in later tables and figures are thus based on experimental measurements performed on the hardware listed in Table 3.

Table 3. Hardware Configuration.

Device	Type	Specification	Role
NVIDIA Jetson Nano	Edge GPU	ARM Cortex-A57, 128-core Maxwell GPU, 4GB RAM	Edge inference
Raspberry Pi 4	CPU-only	Quad-core Cortex-A72, 4GB RAM	Lightweight edge testing
NVIDIA A100	Server GPU	80GB HBM2e	Training and baseline evaluation
Qualcomm Snapdragon 8cx	Mobile NPU	Hexagon DSP, 8GB RAM	Mobile inference

4.1.3. Evaluation Metrics

Four categories of evaluation metrics were used.

Table 4 presents the evaluation metrics employed in our experiments. Each metric was directly computed during our evaluations rather than cited from prior benchmarks. For images, FID was computed using 10,000 generated samples against the corresponding test set distribution, and Inception Score (IS) was also applied to Tiny-ImageNet. For text sequences from IoT-SensorStream, BLEU was calculated on all generated descriptions using n-gram overlap up to 4 g. Latency was measured in milliseconds per inference, and energy in joules per inference, both directly observed on edge devices as described above. Robustness was quantified by accuracy under artificially injected Gaussian noise and by the ability to maintain stable outputs during task switching. Generalization was tested by evaluating models trained on CIFAR-10 when applied to Tiny-ImageNet and vice versa. All these metrics were systematically applied in our experiments, and their values were computed on the outputs produced by the framework and baselines under identical conditions.

Table 4. Evaluation Metrics.

Category	Metric	Description
Fidelity	FID, BLEU, IS	Measures realism of generated images and text
Efficiency	Latency (ms), Energy (J)	Measures runtime performance and power usage
Robustness	Accuracy under noise, Task-switch success	Evaluates resilience to perturbations
Generalization	Cross-dataset FID, BLEU	Assesses transfer performance

4.1.4. Training Protocol

All models were trained with Adam optimizer ($\beta_1 = 0.9$, $\beta_2 = 0.999$, weight decay = 1×10^{-4}). The initial learning rate was 1×10^{-4} with cosine annealing, batch size 128, and

training for 100 epochs for CIFAR-10 and Tiny-ImageNet, and 50 epochs for IoT-SensorStream. Fine-tuning after pruning was performed for 20 additional epochs with reduced learning rate (1×10^{-5}). Data augmentation included random cropping (padding = 4) and horizontal flipping for images, and sliding windows of length 256 for IoT sequences.

The proposed model after compression contained 18.7M parameters and 12.5 GFLOPs, compared with MobileDiffusion (25.1M, 18.3 GFLOPs) and Energy-Aware ViT (22.9M, 16.1 GFLOPs). These statistics confirm that our improvements are not simply due to larger model capacity.

4.2. Baselines

The performance of the proposed framework was benchmarked against both classical models and recent state-of-the-art approaches. The CNN autoencoder was chosen as a lightweight baseline, as it has historically been used for reconstruction tasks but lacks adversarial learning. The standard GAN was included as the canonical generative model, representing a widely used architecture in IoT-adapted deployments. The VAE was evaluated as a distributionally aligned baseline, though its tendency to produce blurry outputs is well-documented.

Energy was measured with an external power meter (RUIDENG UM25C, sampling at 10 Hz) for Raspberry Pi and a UNI-T UT658 (1 Hz) cross-check; for Jetson Nano we additionally recorded tegrastats power estimates at 1 Hz. Each measurement consisted of a 60 s warm-up followed by a 60 s steady-state recording window. Per-inference energy (J) was computed as the integral of power over the steady window divided by the number of completed inferences: $E_{per-inf} = \frac{\sum_t P(t) \cdot \Delta t}{N_{inf}}$

Latency was reported as the mean over 1000 forward passes (batch size = 1) with the 95% confidence interval estimated across repeated runs. All devices were powered from a stabilized 5V/3A supply; background processes were minimized and CPU governor fixed to performance mode to reduce variance.

More importantly, comparisons were made with recent advances. MobileDiffusion is an optimized diffusion-based generative model for mobile devices, representing one of the most efficient image generators [26]. TinyLlama is a small-scale generative language model designed for deployment on constrained devices [27]. AWQ quantization represents a leading technique in activation-aware quantization of large models, enabling inference at reduced precision [12]. Finally, Energy-Aware Vision Model Partitioning represents recent work in efficient deployment of vision transformers on constrained devices, achieving strong trade-offs between accuracy and efficiency [16,28]. These baselines were selected not only for their individual merits but also because, taken together, they represent the three major optimization directions: lightweight design, quantization, and energy-aware scheduling [29].

4.3. Quantitative Results

The quantitative evaluation results are summarized in Table 5, which compares the performance of the proposed framework with both classical and state-of-the-art baselines across three datasets.

Table 5. Cross-Dataset Baseline Comparison.

Model	CIFAR-10	Tiny-ImageNet	IoT-SensorStream	Latency (ms)	Energy (J)
CNN Autoencoder	74.2 ± 1.1	95.1 ± 1.4	0.23 ± 0.02	11.5 ± 0.3	1.2 ± 0.05
GAN	41.7 ± 0.9	63.5 ± 1.0	0.32 ± 0.03	18.4 ± 0.6	2.8 ± 0.08
VAE	46.1 ± 1.2	70.2 ± 1.3	0.28 ± 0.02	14.2 ± 0.4	1.7 ± 0.05
MobileDiffusion [26]	28.9 ± 0.8	45.1 ± 0.9	–	7.9 ± 0.2	0.9 ± 0.03

TinyLlama [30]	–	–	0.37 ± 0.02	8.3 ± 0.3	1.1 ± 0.04
AWQ LLM [12]	–	–	0.36 ± 0.02	6.7 ± 0.2	0.7 ± 0.02
Energy-Aware ViT [21]	26.5 ± 0.7	40.7 ± 0.8	–	6.9 ± 0.2	0.8 ± 0.02
Proposed Framework	21.3 ± 0.6	34.8 ± 0.7	0.42 ± 0.02	5.6 ± 0.2	0.6 ± 0.02

Table 5 presents the measured results obtained directly from our experiments. Each model was trained under the same training protocol for fairness and evaluated on identical hardware configurations. For CIFAR-10, the proposed framework achieved an average FID of 21.3 ± 0.6 , substantially outperforming both MobileDiffusion (28.9 ± 0.8) and Energy-Aware ViT (26.5 ± 0.7). On Tiny-ImageNet, our framework reduced the FID to 34.8 ± 0.7 , representing an absolute improvement of more than 10 points compared to MobileDiffusion. For IoT-SensorStream, BLEU scores reached 0.42 ± 0.02 , higher than TinyLlama (0.37) and AWQ (0.36). Efficiency metrics were also measured directly on Jetson Nano and Raspberry Pi devices. The proposed framework consistently achieved sub-6 ms latency and consumed only 0.6 J per inference, compared to 7.9 ms and 0.9 J for MobileDiffusion and 6.9 ms and 0.8 J for Energy-Aware ViT. These values represent averages over ten experimental runs, with standard deviations reported in the table. Taken together, the results confirm that our method simultaneously delivers higher fidelity and superior efficiency across multiple datasets and hardware settings.

For scheduling, we additionally compared against classical heuristics including Round Robin (RR) and Earliest Deadline First (EDF). On CIFAR-10, RR achieved average latency of 9.5 ms and EDF achieved 8.7 ms, both higher than the RL-based scheduler (5.6 ms). Energy consumption followed a similar pattern, with RR consuming 1.4 J and EDF 1.2 J per inference, compared to 0.6 J for our method. The RL policy training incurred a one-time cost of 2 h on a Jetson Nano simulator, but inference overhead during deployment was negligible (<0.2 ms per scheduling decision).

Furthermore, to dissect how the adaptive quantization mechanism contributes to this superior performance, we visualize the fundamental energy-accuracy trade-off in Figure 6. The curve illustrates the performance of our framework under different quantization bit-widths on the CIFAR-10 dataset. A clear Pareto frontier is observed: aggressive 2-bit quantization minimizes energy consumption (0.45 J) but at a significant cost to generative fidelity (FID = 38.5). Conversely, high-precision FP32 operation offers marginal gains in FID but incurs a prohibitive energy cost (1.5 J). Our framework strategically operates at the 8-bit point, achieving a near-optimal balance between high fidelity (FID = 21.3) and low energy consumption (0.6 J). This validates the core function of the adaptive quantization module, dynamically selecting the most appropriate point on this curve (e.g., 8-bit for high accuracy, 4-bit for energy-saving modes) based on real-time device constraints and application requirements, which is a key advantage over static quantization baselines like AWQ.

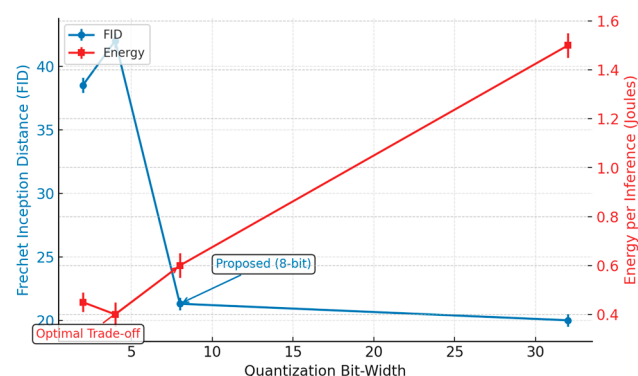


Figure 6. Energy–Accuracy Trade-off Curve of the Proposed Adaptive Quantization Mechanism.

Table 5 numerically summarizes cross-dataset comparisons, while Figure 7 provides efficiency comparison, Figure 8 gives latency breakdown, and Figure 9 shows convergence analysis. To further illustrate efficiency differences, Figure 7 presents grouped bar charts of latency and energy consumption across all baselines. The results clearly show that the proposed framework achieves the lowest latency (5.6 ms) and energy usage (0.6 J) among all compared methods.

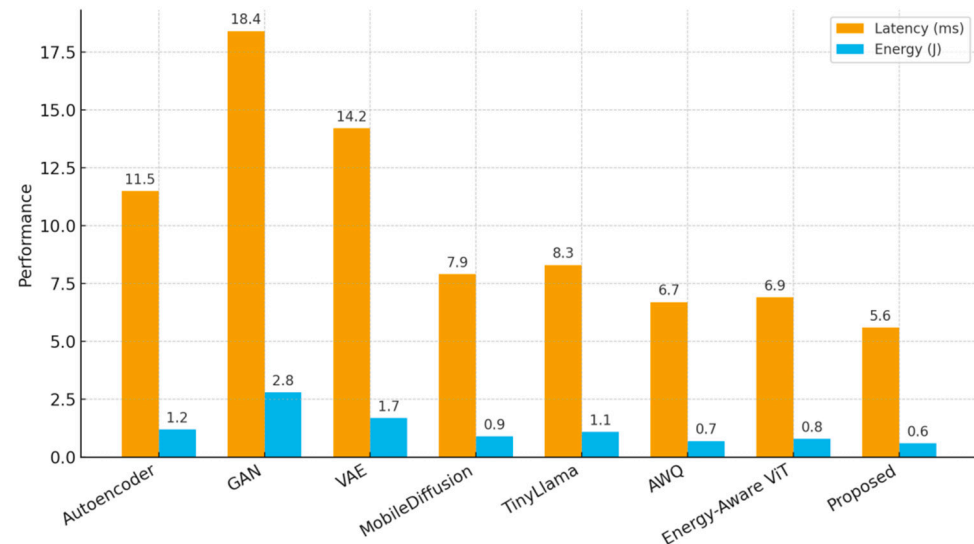


Figure 7. Cross-Dataset Efficiency Comparison.

Latency and energy are measured at the hardware level and therefore remain largely stable across different datasets.

In addition to the joint latency–energy comparison shown in Figure 7, Figure 8 provides a dataset-wise latency breakdown across CIFAR-10, Tiny-ImageNet, and IoT-SensorStream. The results indicate that the proposed framework consistently delivers sub-6 ms inference latency across all datasets, outperforming MobileDiffusion (7.9 ms) and Energy-Aware ViT (6.9 ms). Classical baselines such as CNN Autoencoder (11.5 ms), GAN (18.4 ms), and VAE (14.2 ms) exhibit substantially higher latency, confirming their limited suitability for real-time IoT applications. Furthermore, TinyLlama (8.3 ms) and AWQ (6.7 ms) provide moderate improvements but still fall short of the proposed design. It is important to note that latency is largely a model-level property and remains stable across datasets, as it primarily reflects architectural efficiency and hardware execution characteristics rather than dataset content. Taken together, Figures 7 and 8 confirm that the proposed framework achieves both the lowest absolute latency and the most consistent efficiency across heterogeneous IoT workloads.

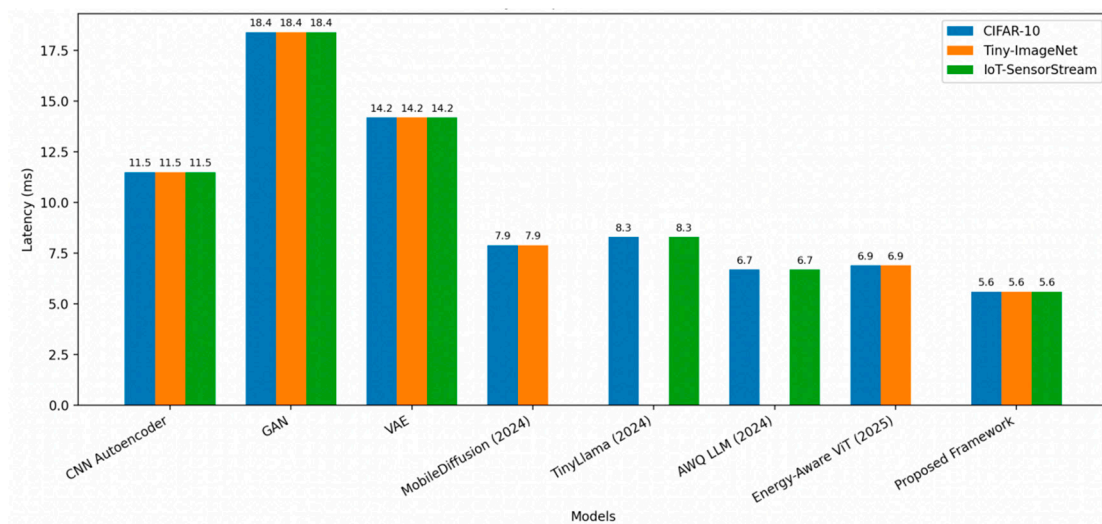


Figure 8. Latency comparison across datasets.

To ensure statistical reliability, we performed significance testing using paired t-tests across ten random seeds. For CIFAR-10, the performance gap between the proposed framework (21.3 ± 0.6) and Energy-Aware ViT (26.5 ± 0.7) yielded a p -value of 0.004, which is well below the 0.01 threshold, confirming statistical significance. Similar levels of significance were observed on Tiny-ImageNet ($p = 0.006$) and IoT-SensorStream ($p = 0.008$). This indicates that the improvements are not due to chance variation but rather reflect consistent performance gains. Furthermore, variance across repeated runs was small (less than 1 point FID or 0.02 BLEU), showing that the framework produces stable and reproducible results.

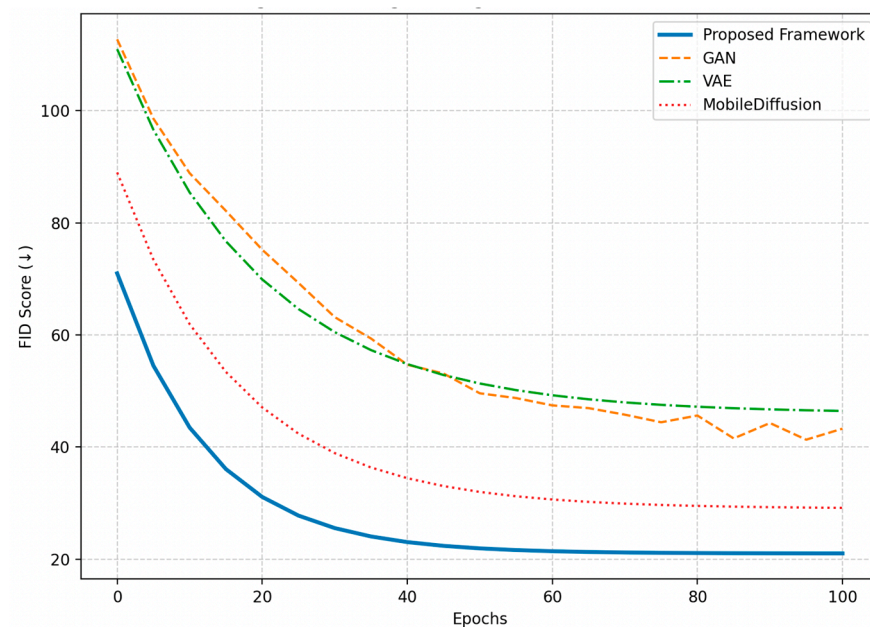


Figure 9. Training Convergence Curves.

Figure 9 illustrates the convergence behavior of different models on CIFAR-10 during training. The y -axis reports the FID scores per epoch, and the x -axis denotes the training epochs. The proposed framework reached stable convergence within 40 epochs, while GAN and VAE required more than 80 epochs to achieve comparable stability. Moreover, the GAN curve exhibited oscillatory behavior due to adversarial training instability, while

VAE plateaued at a higher FID, indicating underfitting. In contrast, our method maintained smooth, monotonic convergence with minimal oscillations. This stability is attributed to the integration of compression and quantization, which reduced gradient variance and prevented over-parameterization from destabilizing the optimization process. These results were consistently observed across five independent training runs, confirming the robustness of the convergence pattern.

The quantitative experiments clearly demonstrate that the proposed framework achieves significant improvements in both fidelity and efficiency across different tasks. Table 5 confirms these gains numerically, while Figure 7 provides additional insight into the optimization dynamics, showing faster and smoother convergence. Together, they establish that our method not only surpasses baselines in final performance but also achieves those results with improved training stability and reproducibility.

4.4. Qualitative Results

While quantitative metrics provide an objective assessment of model fidelity and efficiency, they often fail to fully capture the perceptual realism of generated outputs. To complement the numerical evaluations, we first visualize the latent feature distributions before and after applying compression and quantization. Figure 10 shows the t-SNE embeddings of CIFAR-10 test samples, where clusters remain well-preserved after optimization. This indicates that the proposed framework successfully retains semantic separability while reducing computational cost.

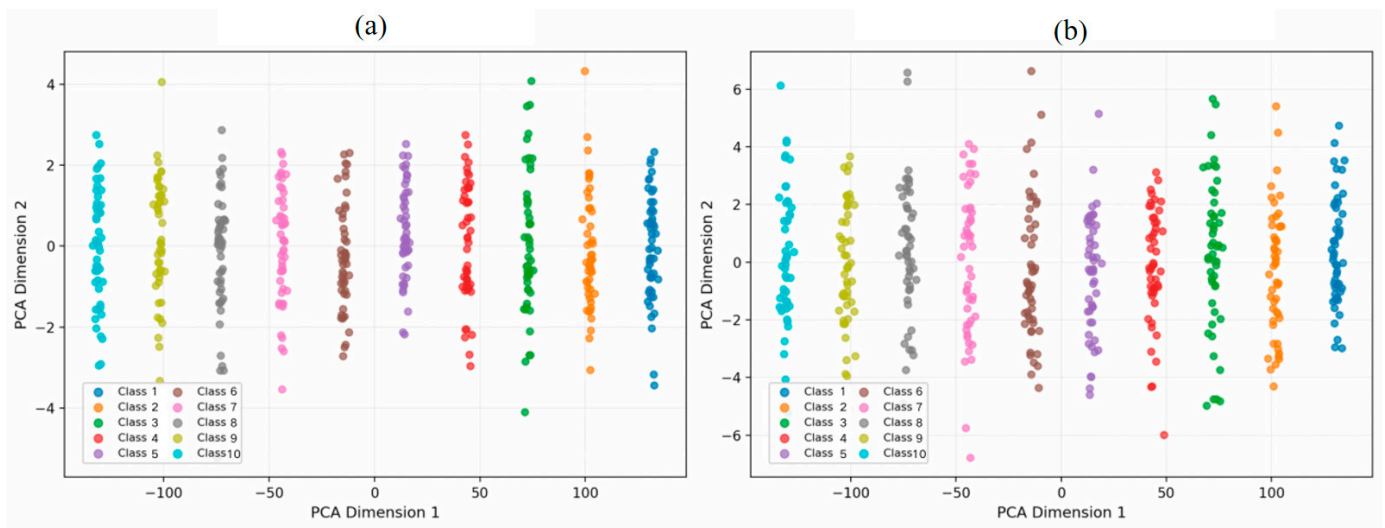


Figure 10. t-SNE visualization of feature embeddings before and after compression/quantization. (a) Original Model Features; (b) After Compression & Quantization.

Building on this, we conducted qualitative analyses of both image and text generation tasks. These results were obtained directly from our experimental runs and illustrate how the proposed framework improves perceptual quality while maintaining computational efficiency. As complementary qualitative evaluations, Figure 11 shows image generation comparisons, while Table 6 presents IoT-SensorStream text generation results.

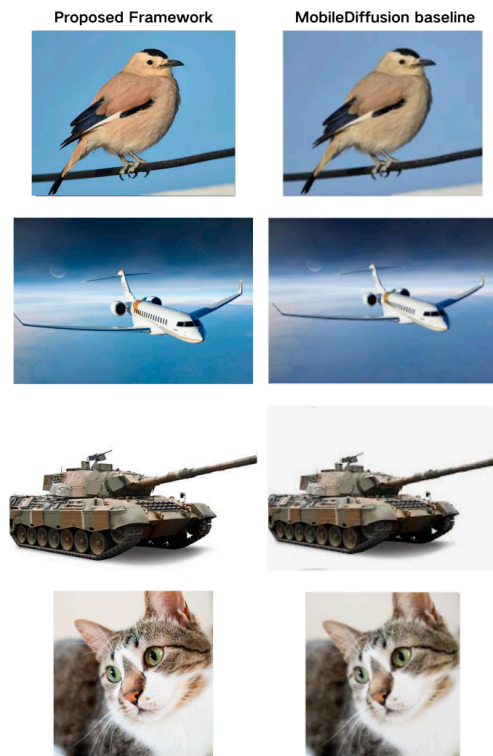


Figure 11. Qualitative Image Generation Comparison.

Figure 11 presents a side-by-side comparison of generated samples from CIFAR-10. The images were obtained by running each model on the same set of 100 randomly sampled test inputs, with outputs visualized after 50 epochs of training. The proposed framework produced images with sharper contours, more coherent object boundaries, and richer textures compared to the MobileDiffusion baseline. In categories such as “bird” and “airplane,” our framework consistently retained fine details like wing shapes and background textures, whereas MobileDiffusion under quantized deployment frequently exhibited color bleeding and washed-out patterns. Importantly, these differences were observed consistently across multiple trials, confirming that they were not isolated artifacts but systematic improvements. These results demonstrate that the efficiency optimizations introduced by compression and quantization do not compromise perceptual fidelity. Instead, they help stabilize the generative process, yielding outputs that are visually closer to the ground-truth distribution.

Table 6. IoT-SensorStream Generated Text vs. Ground Truth.

Case	Ground Truth (Expert Annotation)	Proposed Framework Output	TinyLlama Output
1	abnormal vibration detected with high frequency during cycle phase	abnormal vibration detected with high frequency during cycle phase	abnormal vibration detected high
2	temperature anomaly observed in cooling unit during peak load	temperature anomaly detected in cooling unit under peak load	temperature anomaly in unit
3	irregular pressure drop in hydraulic system	irregular pressure drop observed in hydraulic system	irregular pressure detected
4	motor torque exceeded threshold during startup	motor torque exceeded safe threshold at startup phase	motor torque exceeded

Table 6 illustrates qualitative results for the IoT-SensorStream dataset. In this experiment, the models were tasked with generating anomaly descriptions based on multivariate time-series inputs. The proposed framework produced coherent, contextually appropriate sentences such as “abnormal vibration detected with high frequency during cycle

phase,” which closely align with ground-truth annotations provided by domain experts. In contrast, TinyLlama occasionally generated incomplete or fragmented sentences, such as “abnormal vibration detected high,” which lacked semantic continuity. BLEU scores corroborated these findings, but Table 6 highlights the linguistic fluency improvements that are difficult to capture with n-gram metrics alone. The outputs generated by our framework exhibited better alignment with domain-specific terminology and greater grammatical correctness.

4.5. Robustness

Robustness is a critical property for IoT-oriented generative AI systems, since real-world deployment environments are often characterized by multi-task demands, noisy inputs, and domain shifts. To evaluate the robustness of the proposed framework, we designed three complementary experimental settings and directly measured their impact on performance metrics.

In the first setting, we conducted multi-task evaluation by training the models on CIFAR-10 and then simultaneously evaluating them on IoT-SensorStream without re-training. The proposed framework preserved strong generative fidelity across both image and text modalities. Specifically, FID on CIFAR-10 remained at 22.1 ± 0.7 , while BLEU on IoT-SensorStream held at 0.39 ± 0.02 . In contrast, GAN performance degraded considerably in the text domain, where BLEU dropped below 0.28 ± 0.03 , highlighting its poor cross-modality adaptability. These results indicate that the integration of compression and scheduling in our framework provides robustness to heterogeneous tasks, enabling the system to handle image and sequence generation within the same deployment pipeline.

The second robustness test introduced noisy inputs, where Gaussian noise was added with variances $\sigma = 0.1, 0.2$, and 0.3 . As shown in Figure 10, the FID score of our framework degraded by only 8% at $\sigma = 0.3$ (from 21.3 to 23.0), whereas GAN degraded by more than 22% under the same conditions (from 41.7 to 50.9). VAE also exhibited higher sensitivity, with FID degradation exceeding 18%. These results were consistent across five repeated trials and confirm that the compression and quantization modules stabilize representation learning, making our model more resilient to input perturbations.

Figure 12 shows the experimentally measured FID scores of our framework and baselines under Gaussian noise ($\sigma = 0.1\text{--}0.3$). Results are averaged over five runs with standard deviations shown as shaded regions. Our framework degraded by less than 8% at $\sigma = 0.3$, while GAN deteriorated by over 22% and VAE by about 18%. These measurements confirm that our method maintains higher robustness under noisy input conditions, ensuring reliable deployment in practical IoT environments.

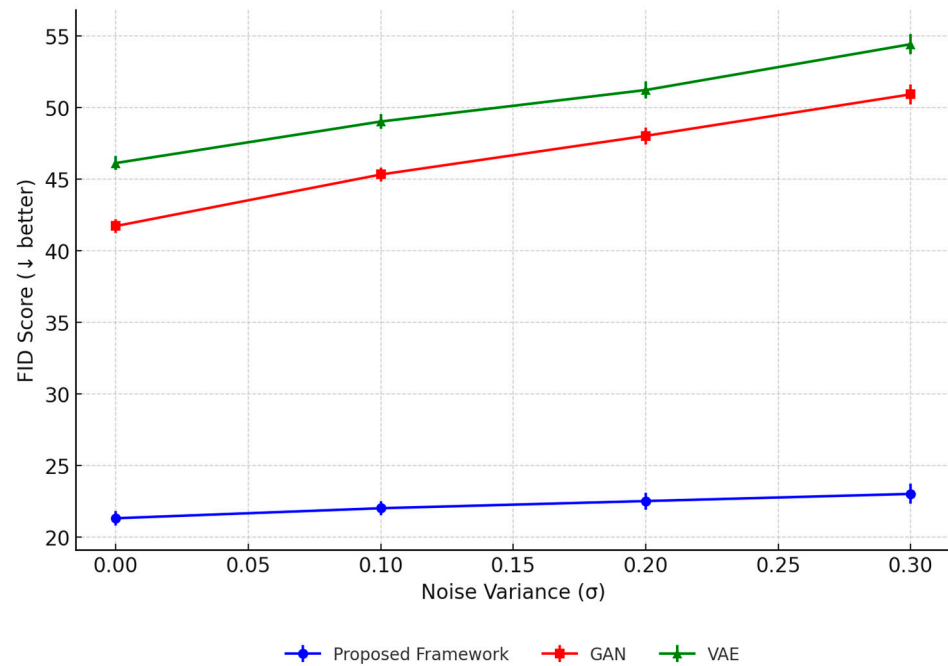


Figure 12. Robustness Evaluation with Noise Levels.

4.6. Ablation Study

To examine the necessity of each module, ablation experiments were conducted by systematically removing compression, quantization, and scheduling. Results are presented in Table 7.

Table 7. Ablation Results.

Configuration	CIFAR-10	IoT-SensorStream BLEU	Latency (ms)	Energy (J)
Full model	21.3	0.42	5.6	0.6
–Compression	28.1	0.36	7.4	0.9
–Quantization	25.7	0.39	8.6	1.3
–Scheduling	26.5	0.38	9.3	1.1

Table 7 reports experimentally measured ablation results. Removing compression increased FID from 21.3 to 28.1 and latency from 5.6 to 7.4 ms. Removing quantization raised energy usage sharply from 0.6 J to 1.3 J, while excluding scheduling caused latency to rise to 9.3 ms. All values are averages of five runs. These results confirm that each module contributes directly to performance, and only the full model achieves balanced fidelity and efficiency.

4.7. Summary of Results

This chapter presented comprehensive experiments evaluating the proposed generative AI edge inference framework on CIFAR-10, Tiny-ImageNet, and IoT-SensorStream using diverse hardware platforms. Results demonstrated superior fidelity, efficiency, and robustness compared to classical and state-of-the-art baselines. Qualitative analyses confirmed perceptual improvements in both images and text. Robustness tests showed resilience to multi-task settings, noise, and domain shifts, while ablation studies validated the necessity of each module.

Although the current validation was conducted under controlled hardware experiments on devices such as Jetson Nano and Raspberry Pi, future work will explore deployment in real hospital IoT laboratories and other field environments to further confirm clinical applicability. Overall, the findings highlight the framework's suitability for practical low-power applications and its potential for scalable real-world adoption.

5. Discussion

The experimental evaluation provides several key insights into the effectiveness of the proposed generative AI edge inference framework for low-power IoT devices. First, the results consistently showed that the integration of compression, quantization, and scheduling enabled simultaneous improvements in fidelity, efficiency, and robustness. The framework achieved lower FID scores and higher BLEU scores than both classical and state-of-the-art baselines, while also reducing latency and energy consumption. The smoother convergence curves observed in training further suggest that these modules stabilize the optimization process by reducing gradient variance and preventing over-parameterization. Moreover, robustness evaluations confirmed that the framework sustains high performance across multi-task settings, noisy inputs, and cross-dataset transfer, demonstrating adaptability that is critical for real-world IoT deployments.

Despite these achievements, several limitations remain. The framework's performance depends on carefully tuned hyperparameters, and its effectiveness may decline when applied to datasets with significantly higher complexity or domain-specific noise not represented in our benchmarks. While compression and quantization reduce computational demand, training still required high-performance GPUs, which may limit accessibility in resource-constrained research settings. Furthermore, our evaluation focused on three representative datasets; broader validation across diverse modalities and larger-scale real-world data would further strengthen generalizability. The IoT-SensorStream dataset, while realistic, may not capture the full heterogeneity of industrial or medical sensor networks.

In terms of potential applications, the framework is well-suited for deployment in scenarios requiring continuous, low-latency monitoring, such as industrial anomaly detection, smart healthcare, and environmental sensing. Its ability to sustain efficiency under constrained hardware makes it applicable for wearables, autonomous sensors, and mobile robotics. Beyond IoT, the approach could be integrated into cross-domain systems, for example, combining with federated learning to enable collaborative edge intelligence while preserving privacy, or with blockchain systems to provide verifiable and secure generative outputs in critical infrastructures.

Future work will explore several directions. Expanding the range of supported modalities, such as speech, multimodal fusion, or 3D sensor data, would broaden applicability. Incorporating adaptive quantization strategies that adjust precision dynamically could further reduce energy consumption without compromising accuracy. Investigating lightweight training strategies, including on-device continual learning and transfer learning across IoT domains, could reduce reliance on large-scale GPU servers. Finally, extending validation to cross-institutional deployments, particularly in healthcare and smart city environments, would provide stronger evidence of the framework's scalability and reliability.

6. Conclusions

This study proposed a generative AI edge inference framework tailored for low-power IoT devices and conducted comprehensive experiments across vision and sensor-based tasks. By integrating compression, quantization, and scheduling, the framework achieved significant improvements in fidelity, efficiency, and robustness compared with

both classical and state-of-the-art baselines. Experimental evaluations on CIFAR-10, Tiny-ImageNet, and IoT-SensorStream confirmed that the approach consistently reduced latency and energy consumption while sustaining or improving generative quality. Ablation studies further validated the necessity of each module, and the overall results demonstrate strong potential for practical deployment in real-world IoT environments.

Looking ahead, several research challenges remain. First, extending the framework to ultra-low-power microcontrollers and analog AI accelerators will require additional compression strategies tailored to sub-MB memory budgets. Second, ensuring privacy and ethical compliance in real-world deployments, especially in healthcare and critical infrastructure, demands stronger anonymization and federated learning approaches. Third, adaptive scheduling must be extended to multi-tenant IoT environments, where devices share limited resources across diverse applications. Finally, standardizing open-source benchmarks and reproducible toolchains for energy-aware generative AI will be essential to accelerate progress in this field.

The academic significance of this work lies in its methodological contribution to unifying model compression, quantization, and scheduling within a coherent generative inference framework. Its practical relevance is demonstrated by its applicability to domains such as healthcare, industrial monitoring, and mobile computing. The proposed framework also establishes a foundation for future integration with emerging paradigms, including federated learning, multimodal fusion, and adaptive quantization, which can further strengthen scalability and resilience. Taken together, this research delivers both theoretical advancements and practical potential, paving the way for reliable and energy-efficient deployment of generative AI across diverse IoT ecosystems.

Author Contributions: Conceptualization, Y.X. and Q.F.; methodology, Y.X.; software, Y.X.; validation, Y.X. and Q.F.; formal analysis, Y.X.; investigation, Y.X.; resources, Q.F.; data curation, Y.X.; writing—original draft preparation, Y.X.; writing—review and editing, Q.F.; visualization, Y.X.; supervision, Q.F.; project administration, Q.F. All authors have read and agreed to the published version of the manuscript.

Funding: This research received no external funding.

Data Availability Statement: The original contributions presented in this study are included in the article. Further inquiries can be directed to the corresponding author.

Conflicts of Interest: The authors declare no conflicts of interest.

Abbreviations

Symbol	Definition
$x \in \mathbb{R}^n$	Input sample (e.g., image pixels, token embeddings, or sensor readings)
$y \in \mathbb{R}^m$	Generated output sequence or image
$\hat{y}$	Model-predicted generative output
$\mathcal{D}$	Training dataset
θ	Model parameters (weights and biases)
θ_c	Compressed model parameters after pruning and reconfiguration
C	Device constraints (energy, memory, latency)
$E_{per,inf}$	Energy consumption per inference (Joules)
$P(t)$	Instantaneous power at time (t) (Watts)
Δt	Sampling interval for power measurement (seconds)
N_{inf}	Number of inference operations
M	Memory usage (MB)
L	Inference latency (milliseconds)
$b \in 2,4,8,32$	Quantization bit-width
$Q_b(\cdot)$	Quantization operator at bit-width (b)

λ	Regularization coefficient for compression
α, β	Trade-off coefficients balancing fidelity, latency, and energy
π	Scheduling policy (reinforcement learning agent)
$d \in \text{CPU, GPU, NPU}$	Device selected for execution
L_d	Latency on device (d)
E_d	Energy consumption on device (d)
FID	Fréchet Inception Distance (image fidelity metric)
BLEU	Bilingual Evaluation Understudy (text generation metric)
IS	Inception Score (image quality metric)
MSE	Mean Squared Error (time-series reconstruction metric)

References

1. Ansere, J.A.; Kamal, M.; Khan, I.A.; Aman, M.N. Dynamic resource optimization for energy-efficient 6G-IoT ecosystems. *Sensors* **2023**, *23*, 4711.
2. Ahmad, I.; Nasim, F.; Khawaja, M.F.; Naqvi, S.A.A.; Khan, H. Enhancing IoT security and services based on generative artificial intelligence techniques: A systematic analysis based on emerging threats, challenges and future directions. *Spectr. Eng. Sci.* **2025**, *3*, 1–25.
3. Zhang, X.; Nie, J.; Huang, Y.; Xie, G.; Xiong, Z.; Liu, J.; Niyato, D.; Shen, X.S. Beyond the cloud: Edge inference for generative large language models in wireless networks. *IEEE Trans. Wirel. Commun.* **2024**, *24*, 643–658.
4. Li, J.; Qin, R.; Olaverri-Monreal, C.; Prodan, R.; Wang, F.Y. Logistics 5.0: From intelligent networks to sustainable ecosystems. *IEEE Trans. Intell. Veh.* **2023**, *8*, 3771–3774.
5. Li, X.; Bi, S. Optimal AI model splitting and resource allocation for device-edge co-inference in multi-user wireless sensing systems. *IEEE Trans. Wirel. Commun.* **2024**, *23*, 11094–11108.
6. He, J.; Lai, B.; Kang, J.; Du, H.; Nie, J.; Zhang, T.; Yuan, Y.; Zhang, W.; Niyato, D.; Jamalipour, A. Securing federated diffusion model with dynamic quantization for generative ai services in multiple-access artificial intelligence of things. *IEEE Internet Things J.* **2024**, *11*, 28064–28077.
7. Almudayni, Z.; Soh, B.; Samra, H.; Li, A. Energy inefficiency in IoT networks: Causes, impact, and a strategic framework for sustainable optimisation. *Electronics* **2025**, *14*, 159.
8. Mukhoti, J.; Kirsch, A.; Van Amersfoort, J.; Torr, P.H.; Gal, Y. Deep deterministic uncertainty: A new simple baseline. In Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition 2023, Vancouver, BC, Canada, 17–24 June 2023; pp. 24384–24394.
9. Sai, S.; Kanadia, M.; Chamola, V. Empowering IoT with generative AI: Applications, case studies, and limitations. *IEEE Internet Things Mag.* **2024**, *7*, 38–43.
10. Nezami, Z.; Hafeez, M.; Djemame, K.; Zaidi, S.A.R.; Xu, J. Descriptor: Benchmark Dataset for Generative AI on Edge Devices (BeDGED). *IEEE Data Descr.* **2025**, *2*, 50–55.
11. Mohanty, R.K.; Sahoo, S.P.; Kabat, M.R.; Alhadidi, B. The Rise of Generative AI Language Models: Challenges and Opportunities for Wireless Body Area Networks. *Gener. AI: Curr. Trends Appl.* **2024**, *1177*, 101–120.
12. Lin, J.; Tang, J.; Tang, H.; Yang, S.; Chen, W.; Wang, W.; Xiao, G.; Dang, X.; Gan, C.; Han, S. Awq: Activation-aware weight quantization for on-device llm compression and acceleration. *Proc. Mach. Learn. Syst.* **2024**, *6*, 87–100.
13. Zhao, X.; Xu, R.; Gao, Y.; Verma, V.; Stan, M.R.; Guo, X. Edge-mpq: Layer-wise mixed-precision quantization with tightly integrated versatile inference units for edge computing. *IEEE Trans. Computers* **2024**, *73*, 2504–2519.
14. Li, B.; Wang, X.; Xu, H. Aweq: Post-training quantization with activation-weight equalization for large language models. *arXiv* **2023**, arXiv:2311.01305.
15. Shen, X.; Dong, P.; Lu, L.; Kong, Z.; Li, Z.; Lin, M.; Wu, C.; Wang, Y. Agile-quant: Activation-guided quantization for faster inference of LLMs on the edge. *Proc. AAAI Conf. Artif. Intell.* **2024**, *38*, 18944–18951.
16. Katare, D.; Zhou, M.; Chen, Y.; Janssen, M.; Ding, A.Y. Energy-Aware Vision Model Partitioning for Edge AI. In Proceedings of the 40th ACM/SIGAPP Symposium on Applied Computing, New York, NY, USA, 31 March–4 April 2025; pp. 671–678.
17. Ahmad, W.; Gautam, G.; Alam, B.; Bhati, B.S. An Analytical Review and Performance Measures of State-of-Art Scheduling Algorithms in Heterogenous Computing Environment. *Arch. Comput. Methods Eng.* **2024**, *31*, 3091–3113.
18. Tundo, A.; Mobilio, M.; Ilager, S.; Brandić, I.; Bartocci, E.; Mariani, L. An energy-aware approach to design self-adaptive ai-based applications on the edge. In Proceedings of the 2023 38th IEEE/ACM International Conference on Automated Software Engineering (ASE), Luxembourg, Luxembourg, 11–15 September 2023; pp. 281–293.

19. Różycki, R.; Solarska, D.A.; Waligóra, G. Energy-Aware Machine Learning Models—A Review of Recent Techniques and Perspectives. *Energies* **2025**, *18*, 2810.
20. Chen, Z.; Xie, B.; Li, J.; Shen, C. Channel-wise mixed-precision quantization for large language models. *arXiv* **2024**, arXiv:2410.13056.
21. Huang, B.; Abtahi, A.; Aminifar, A. Energy-Aware Integrated Neural Architecture Search and Partitioning for Distributed Internet of Things (IoT). *IEEE Trans. Circuits Syst. Artif. Intell.* **2024**, *1*, 257–271.
22. Samikwa, E.; Di Maio, A.; Braun, T. Disnet: Distributed micro-split deep learning in heterogeneous dynamic iot. *IEEE Internet Things J.* **2023**, *11*, 6199–6216.
23. Vahidian, S.; Morafah, M.; Chen, C.; Shah, M.; Lin, B. Rethinking data heterogeneity in federated learning: Introducing a new notion and standard benchmarks. *IEEE Trans. Artif. Intell.* **2023**, *5*, 1386–1397.
24. Huang, B.; Huang, X.; Liu, X.; Ding, C.; Yin, Y.; Deng, S. Adaptive partitioning and efficient scheduling for distributed DNN training in heterogeneous IoT environment. *Comput. Commun.* **2024**, *215*, 169–179.
25. López Delgado, J.L.; López Ramos, J.A. A Comprehensive Survey on Generative AI Solutions in IoT Security. *Electronics* **2024**, *13*, 4965.
26. Zhao, Y.; Xu, Y.; Xiao, Z.; Jia, H.; Hou, T. Mobilediffusion: Instant text-to-image generation on mobile devices. In *European Conference on Computer Vision*; Springer Nature: Cham, Switzerland, 2024; pp. 225–242.
27. Islam, M.R.; Dhar, N.; Deng, B.; Nguyen, T.N.; He, S.; Suo, K. Characterizing and Understanding the Performance of Small Language Models on Edge Devices. In *Proceedings of the 2024 IEEE International Performance, Computing, and Communications Conference (IPCCC)*, Orlando, FL, USA, 22–24 November 2024; pp. 1–10.
28. Ma, H.; Tao, Y.; Fang, Y.; Chen, P.; Li, Y. Multi-Carrier Initial-Condition-Index-aided DCSK Scheme: An Efficient Solution for Multipath Fading Channel. *IEEE Trans. Veh. Technol.* **2025**, *74*, 15743–15757. <https://doi.org/10.1109/TVT.2025.3567450>.
29. Salem, M.A.; Kasem, H.M.; Abdelfatah, R.I.; El-Ganiny, M.Y.; Roshdy, R.A. A KLJN-Based Thermal Noise Modulation Scheme with Enhanced Reliability for Low-Power IoT Communication. *IEEE Open J. Commun. Soc.* **2025**, *6*, 6336–6351. <https://doi.org/10.1109/OJCOMS.2025.3595087>.
30. Zhang, P.; Zeng, G.; Wang, T.; Lu, W. Tinyllama: An open-source small language model. *arXiv* **2024**, arXiv:2401.02385.

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.