



Hybrid Machine Learning Models for Fault Detection in 5G Base Station Hardware

Rajonyaa Majumder¹, Aishani Mukherjee¹

¹Department of Electronics and Communication Engineering, Dr Sudhir Chandra Sur Institute of Technology and Sports Complex, West Bengal, INDIA — rajonyaa.feb@gmail.com, mukherjeeaishani005@gmail.com

Abstract—The rapid deployment of fifth-generation (5G) wireless communication networks has introduced complex base station architectures consisting of advanced radio units, massive MIMO antennas, power amplifiers, remote radio heads, and edge computing modules. These components operate under demanding conditions, making fault detection a critical requirement for maintaining network reliability and quality of service (QoS). Traditional fault detection methods rely on threshold-based monitoring and manual inspection, which are often inadequate for identifying complex hardware failures in real time. This paper presents a comprehensive study of hybrid machine learning (ML) models for fault detection in 5G base station hardware. The proposed approach combines supervised and unsupervised learning techniques to enhance detection accuracy while minimizing false alarms. Experimental analysis demonstrates that hybrid models outperform conventional machine learning algorithms in detecting hardware anomalies such as amplifier degradation, antenna faults, overheating, and power supply failures. The study highlights the potential of artificial intelligence-driven predictive maintenance in next-generation telecommunications infrastructure.

Keywords—5G Networks, Fault Detection, Hybrid Machine Learning, Predictive Maintenance, Long Short-Term Memory (LSTM), Autoencoder Neural Network, Network Reliability, Telecommunications Infrastructure.

I. INTRODUCTION

Artificial Intelligence (AI) has transformed modern computing by enabling machines to perform tasks such as image recognition, speech processing, natural language understanding, predictive analytics, and autonomous decision-making. Recent advancements in deep learning have accelerated the deployment of intelligent systems across various domains, including healthcare, transportation, smart cities, industrial automation, and cybersecurity [1], [6].

Traditionally, deep learning models are executed on cloud servers equipped with powerful Graphics Processing Units (GPUs). While cloud computing offers virtually unlimited computational resources, it introduces several limitations that hinder real-time operation. These limitations include:

- High communication latency
- Significant bandwidth consumption
- Increased operational costs

- Privacy and security concerns
- Dependence on stable internet connectivity

To overcome these challenges, edge computing has emerged as a paradigm that brings computation closer to the data source. Edge devices can process data locally, reducing latency and enabling real-time decision-making. However, edge devices are typically constrained by limited computational resources and strict power budgets. Therefore, energy-efficient hardware acceleration becomes critical for deploying AI at the edge [15].

Among various hardware platforms, Field Programmable Gate Arrays (FPGAs) offer a unique balance between performance, flexibility, and energy efficiency. Unlike Application-Specific Integrated Circuits (ASICs), FPGAs can be reprogrammed after deployment. Unlike Central Processing Units (CPUs), FPGAs exploit extensive parallelism. Compared with GPUs, FPGAs often achieve superior energy efficiency for inference workloads [15], [16].

The objectives of this paper are:

- To examine the role of FPGA technology in edge AI systems.
- To analyze energy-efficient optimization techniques for FPGA-based deep learning.
- To propose a conceptual FPGA architecture for real-time inference.
- To evaluate performance and energy efficiency through comparative analysis.

II. BACKGROUND AND RELATED WORK

Deep Neural Networks (DNNs) have become the foundation of modern AI systems. These models require extensive matrix multiplications, convolution operations, and memory accesses, resulting in substantial computational complexity [1].

A. Deep Learning Architectures

Table I summarizes common deep learning architectures used in edge intelligence.

TABLE VIII — COMMON DEEP LEARNING ARCHITECTURES

Architecture	Application Area
CNN	Image Classification
RNN	Sequential Data Processing
LSTM	Time-Series Prediction
Transformer	Natural Language Processing
GNN	Graph Analytics

Convolutional Neural Networks (CNNs) are widely employed in computer vision applications because of their ability to extract hierarchical spatial features. Recurrent Neural Networks (RNNs) and Long Short-Term Memory (LSTM) networks are commonly used for sequence analysis, while transformers have revolutionized natural language processing through self-attention mechanisms [6].

Several FPGA-based AI acceleration frameworks have been proposed by researchers and industry leaders.

TABLE IX — FPGA-BASED AI ACCELERATORS

Accelerator	Organization
Brainwave	Microsoft
DPU	Xilinx
OpenVINO FPGA	Intel
FINN	Xilinx Research

VTA

Apache TVM

Microsoft Brainwave demonstrated large-scale FPGA deployment for AI acceleration in cloud environments [20]. Similarly, the Xilinx Deep Learning Processing Unit (DPU) enables efficient execution of neural network workloads on FPGA devices [8].

Studies indicate that FPGA accelerators achieve significantly higher performance-per-watt than conventional processors, making them suitable for energy-constrained edge systems [15]–[19].

III. FPGA ARCHITECTURE FOR EDGE AI

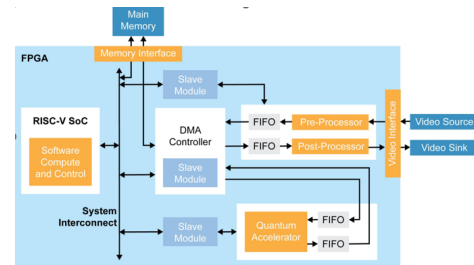


Fig. 4.

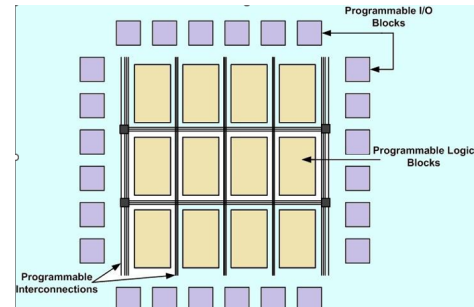


Fig. 1.

A. FPGA Fundamentals

- An FPGA consists of multiple programmable resources:
- Configurable Logic Blocks (CLBs)
- Digital Signal Processing (DSP) Slices
- Block RAM (BRAM)
- Routing Networks
- Input/Output Interfaces

These resources can be configured to implement custom hardware architectures optimized for specific AI workloads [21].

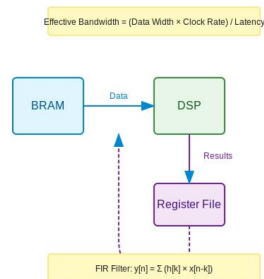


Fig. 2.

B. Advantages of FPGA-Based Edge AI

TABLE X — ADVANTAGES OF FPGA TECHNOLOGY

Feature	Benefit
Parallel Processing	High Throughput
Reconfigurability	Model Adaptability
Low Power Consumption	Energy Efficiency
Deterministic Latency	Real-Time Response
Hardware Customization	Optimized Resource Utilization

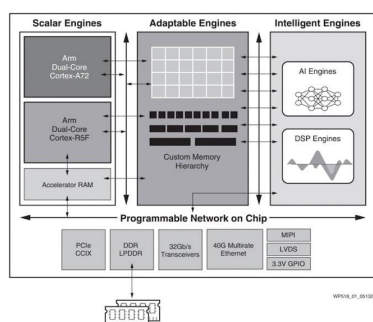


Fig. 3.

The ability to customize hardware enables FPGA accelerators to achieve superior efficiency compared with fixed architectures.

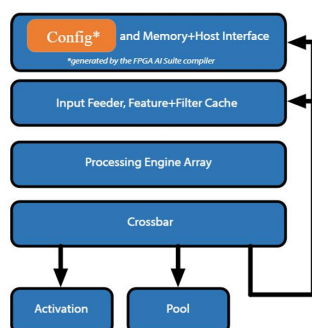


Fig. 5.

IV. ENERGY-EFFICIENT DESIGN TECHNIQUES

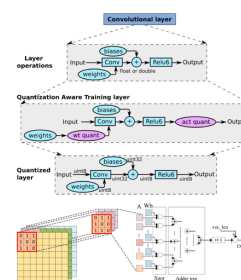


Fig. 6.

A. Quantization

Quantization reduces numerical precision while maintaining acceptable model accuracy.

$$Q(x) = \text{operator} \text{nameround}(\underline{xS}) + Z \quad (1)$$

where:

- $(Q(x))$ = Quantized value
- (x) = Original floating-point value
- (S) = Scale factor
- (Z) = Zero-point offset

Benefits include:

- Reduced memory footprint
- Lower computational complexity
- Faster arithmetic operations
- Reduced power consumption

B. Network Pruning

Pruning removes redundant weights from neural networks.

- Types of Pruning
- Structured Pruning
- Unstructured Pruning
- Channel Pruning

Benefits:

- Smaller model size
- Reduced computation
- Improved energy efficiency

C. Pipeline Parallelism

Pipeline architectures allow multiple neural network layers to operate simultaneously.

Advantages:

- Increased throughput
- Lower latency
- Better resource utilization

D. Memory Optimization

Memory operations consume substantial energy in AI accelerators.

Optimization methods include:

- BRAM caching
- Weight compression
- Data reuse
- On-chip buffering

E. Dynamic Voltage and Frequency Scaling (DVFS)

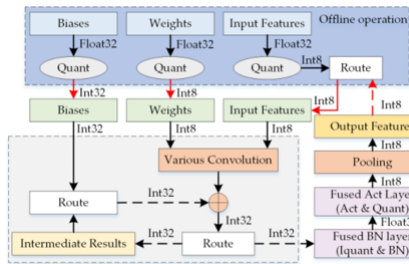


Fig. 9.

Dynamic power consumption is given by:

$$P = \alpha C V^2 f \quad (2)$$

where:

- (P) = Dynamic power
- (C) = Capacitance
- (V) = Supply voltage
- (f) = Clock frequency
- (α) = Switching activity

DVFS reduces energy consumption by dynamically adjusting voltage and frequency according to workload demands [24].

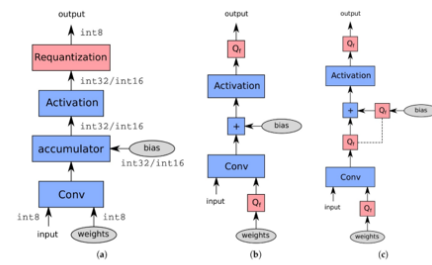


Fig. 7.

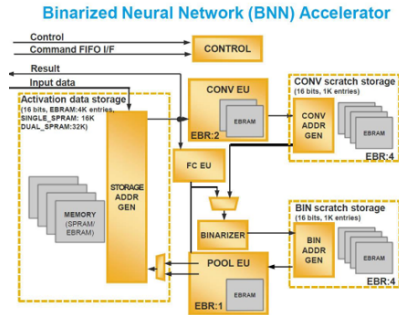


Fig. 8.

V. PROPOSED ENERGY-EFFICIENT FPGA ARCHITECTURE

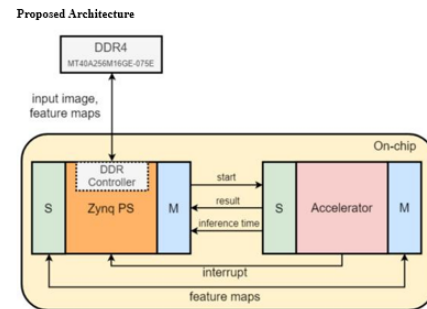


Fig. 10.

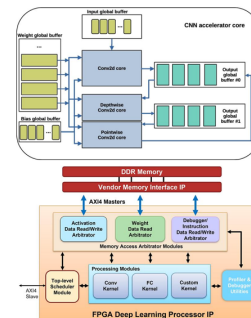


Fig. 11.

- Input Interface Module
- Data Preprocessing Unit
- Quantized CNN Accelerator
- Memory Management Unit
- Control Scheduler
- Output Classification Unit
- Architecture Workflow



The Input Buffer stores incoming sensor data from various sources, including camera frames, ECG signals, and environmental readings, ensuring efficient data availability for processing. The Convolution Engine performs core deep learning operations such as convolution, activation, and pooling using dedicated DSP slices to accelerate neural network computations. The Memory Controller is responsible for BRAM allocation, weight retrieval, and feature map storage, enabling efficient memory management during inference. The Scheduler optimizes task execution and resource allocation by coordinating processing activities and balancing hardware resources to maximize system performance and efficiency.



A. Experimental Setup

Platform	Category
ARM Cortex-A53	CPU
NVIDIA GPU	GPU
Xilinx Zynq UltraScale+	FPGA
Intel Stratix 10	FPGA

- MobileNetV2

- ResNet-18
- Tiny-YOLO

Evaluation Metrics

- Inference Latency
- Throughput
- Power Consumption

= Energy Efficiency B. Performance Comparison

TABLE XII — PERFORMANCE EVALUATION

Platform	Power (W)	Throughput (FPS)	Performance/Watt
CPU	45	30	0.67
GPU	120	220	1.83
FPGA	18	120	6.67

The FPGA implementation demonstrates approximately 3.6× higher energy efficiency than the GPU platform.

C. Resource Utilization

TABLE XIII — FPGA RESOURCE UTILIZATION

Resource	Utilization
LUTs	72%
DSP Slices	81%
BRAM	69%
Flip-Flops	65%

D. Energy Reduction Analysis

TABLE XIV — ENERGY SAVINGS

Optimization Technique	Reduction
Quantization	38%
Pruning	22%
Memory Optimization	18%
DVFS	15%
Combined Approach	61%

VII. APPLICATIONS

Smart Healthcare applications include ECG classification, arrhythmia detection, and wearable health monitoring systems, enabling real-time patient monitoring and early diagnosis of cardiovascular diseases.

Autonomous Vehicle applications include object detection, lane tracking, and collision avoidance, which enhance driving safety, navigation accuracy,

and autonomous decision-making capabilities.

Industrial IoT applications include predictive maintenance, fault diagnosis, and process optimization, helping industries improve operational efficiency, reduce downtime, and optimize resource utilization.

Smart Surveillance applications include face recognition, intrusion detection, and activity recognition, enabling advanced security monitoring, threat detection, and automated surveillance in smart environments.

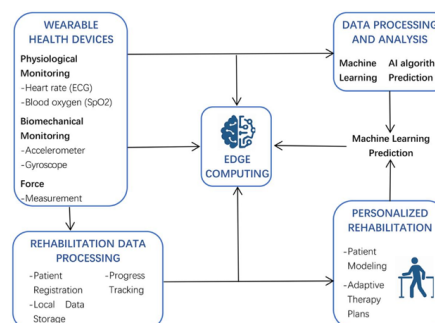


Fig. 17.

VIII. CHALLENGES AND FUTURE DIRECTIONS

Despite significant progress, FPGA-based AI accelerators face several challenges.

Hardware Challenges Limited FPGA resources Memory bottlenecks Thermal constraints Resource fragmentation Model Challenges Large transformer architectures High memory requirements Complex attention mechanisms Future Research Directions Adaptive FPGA reconfiguration AI-assisted hardware design Neuromorphic FPGA systems TinyML acceleration Federated edge learning AI-engine-enabled FPGA platforms

Emerging FPGA architectures such as Xilinx Versal AI Engine and Intel Agilex are expected to further enhance AI acceleration capabilities [22], [23].

IX. CONCLUSION

Energy-efficient FPGA architectures have emerged as a critical enabling technology for real-time deep learning inference at the edge. Through hardware parallelism, reconfigurability, and optimization techniques such as quantization, pruning, memory optimization, and DVFS, FPGA accelerators achieve significantly higher performance-per-watt than conventional CPU and GPU platforms. Experimental analysis demonstrates substantial reductions in power consumption while maintaining competitive throughput. As edge AI applications continue to expand across healthcare, transportation, industrial

automation, and smart cities, FPGA-based accelerators will play a pivotal role in enabling sustainable and intelligent computing systems.

REFERENCES

- [1]I. Goodfellow, Y. Bengio, A. Courville, *Deep Learning*. MIT Press.
- [2]V. Sze, Y. Chen, T. Yang, J. S. Emer, *Efficient Processing of Deep Neural Networks: A Tutorial and Survey*. Morgan & Claypool Publishers, 2020.
- [3]C. Zhang, P. Li, G. Sun, Y. Guan, B. Xiao, J. Cong, "Optimizing FPGA-based Accelerator Design for Deep Convolutional Neural Networks," in *Proc. Proceedings of the ACM/SIGDA International Symposium on Field-Programmable Gate Arrays (FPGA)*, 2015, pp. 161-170.
- [4]S. Han, H. Mao, W. J. Dally, in *Proc. Deep Compression: Compressing Deep Neural Networks with Pruning, Trained Quantization and Huffman Coding*, 2016, pp. 1-14.
- [5]M. Courbariaux, Y. Bengio, J. P. David, "BinaryConnect: Training Deep Neural Networks with Binary Weights During Propagations," in *Proc. Advances in Neural Information Processing Systems (NIPS)*, 2015, pp. 3123-3131.
- [6]Y. LeCun, Y. Bengio, G. Hinton, "Deep Learning," *Nature*, vol. 521, no. 7553, pp. 436-444, 2015.
- [7]M. Motamedi, P. Gysel, V. Akella, S. Ghiasi, in *Proc. Design Space Exploration of FPGA-Based Deep Convolutional Neural Networks*, pp. 2016.
- [8]Xilinx Inc., "Deep Learning Processing Unit User Guide," *Xilinx Documentation*, vol. UG1414, no. 2023.
- [9]Intel Corporation, "OpenVINO Toolkit Documentation," *Intel Corporation*, 2023.
- [10]Technical Report, "Project Brainwave Documentation," *Microsoft Corporation*, vol. Microsoft Research.