



Adaptive Channel Estimation of Beamforming O+F OFDM Signal in 6G Optical Wireless Communication Systems using Deep Machine Learning Techniques with FPGA Implementation

Fatai Ahmed-Ade¹ (<https://orcid.org/0009-0003-4123-7338>), Vincent Andrew Akpan²,
(<https://orcid.org/0000-0002-5155-2863>), Margret Adebimpe Umeche³,
(<https://orcid.org/0009-0003-0113-0325>), Daniel Yusuf Odoma⁴ (<https://orcid.org/0009-0009-8578-103X>), Ninma Mary Esuga-Mopah⁵ (<https://orcid.org/0009-0002-3534-3428>)

¹Department of Physics, Federal University Lokoja, Lokoja, Kogi State, Nigeria

²Department of Biomedical Engineering, The Federal University of Technology, Akure, Ondo State, Nigeria

³Department of Physics, Kogi State University, Kabba, Kogi State, Nigeria

⁴Department of Physics Electronics, The Federal Polytechnic, Idah, Kogi State, Nigeria

⁵Department of Science Laboratory Technology, Kogi State Polytechnic, Lokoja, Kogi State, Nigeria

Corresponding Author: Vincent Andrew Akpan

ABSTRACT: The rapid evolution toward sixth-generation (6G) networks demands ultra-high data rates, sub-millisecond latency, and exceptional reliability, particularly in optical wireless communication (OWC) systems. This paper presents a comprehensive hardware-accelerated framework that integrates Optical Plus Filtered Orthogonal Frequency Division Multiplexing (O+F OFDM), a Temporal Convolutional Network (TCN)-based adaptive channel predictor, and an intelligent beamforming engine, all implemented on a Xilinx Virtex-7 FPGA. The proposed O+F OFDM modulator effectively mitigates critical limitations of conventional optical OFDM by achieving a 4.1 dB reduction in peak-to-average power ratio (from 11.2 dB to 7.1 dB) and 18 dB suppression of out-of-band emissions at 250 kHz offset. At the receiver, the system delivers a 2.3 dB SNR gain at a BER of 10^{-4} and improves power efficiency by 26%. For real-time channel estimation under dynamic conditions, a TCN-based predictor is developed and rigorously benchmarked against Long Short-Term Memory (LSTM) networks. The TCN demonstrates superior hardware suitability, achieving inference latency of 1.2 ms on the FPGA compared to 5.8 ms for LSTM in simulation. It also requires significantly lower memory (2.74 MB vs. 6.53 MB in INT8) while providing better prediction accuracy (indoor MSE of 0.028 vs. 0.050). Overall, the TCN attains 94.3% channel prediction accuracy at a 50 ms lookahead horizon, reducing beam acquisition time by 67%. The complete FPGA implementation achieves an end-to-end latency of 4.7 μ s for the signal pipeline and 3.7 ms for the full system (well within the 5 ms URLLC requirement), with a total power consumption of 8.2 W representing a 44–47% reduction compared to equivalent DSP implementations. Real-world validation in Lagos, Nigeria, under challenging atmospheric conditions (high solar irradiance and harmattan dust) confirmed over 97% link availability and 40% higher throughput than the baseline. These results establish the proposed TCN-enhanced O+F OFDM framework as a practical, energy-efficient, and deployable solution for 6G optical wireless systems.

KEYWORDS: Sixth Generation (6G) Networks, Beamforming, Adaptive Channel Estimation, Deep Neural Networks, FPGA Implementation, Optical Wireless Communication, O+F OFDM

Received 04 Aug., 2026; Revised 14 Aug., 2026; Accepted 16 Aug., 2026 © The author(s) 2026.

Published with open access at www.questjournals.org

I. INTRODUCTION

The global telecommunications landscape is undergoing a fundamental paradigm shift as the research community transitions from fifth-generation (5G) networks toward the more ambitious architecture of sixth-generation (6G) systems, broadly anticipated for commercial deployment by 2030. Although 5G has demonstrably advanced the state of enhanced mobile broadband with measured peak data rates approaching 10 Gbps under favourable propagation conditions the technical and socioeconomic demands projected for the coming decade require capabilities that exceed 5G specifications by several orders of magnitude. Sixth-

generation networks are envisioned to deliver terabit-per-second aggregate data rates, sub-millisecond end-to-end latency, reliability approaching 99.9999%, and device densities reaching 10 million nodes per square kilometer [1, 2]. Meeting these targets cannot be achieved through incremental improvements to existing radio-frequency physical layer technologies; they demand a fundamental reconceptualization of how spectrum, spatial resources, and signal processing intelligence are co-designed and jointly optimized at every layer of the protocol stack.

Beamforming is a spatial filtering technique used to enhance and improve the signal-to-noise (SNR) of received signals by eliminating undesirable interference sources and focus on the transmitted and received signal. Beamforming is a signal-processing technique used in antenna arrays to direct wireless signals (such as signal processing, radar, biomedical and particularly in communications including Wi-Fi and 5G) precisely toward a specific device rather than broadcasting in all directions [3, 4]. By controlling signal phase and amplitude, it creates constructive interference, strengthening the signal while reducing interference, resulting in faster speeds and better range.

Beamforming in 6G Optical Wireless Communication Systems (OWCS) utilizes precise, agile light manipulation to overcome atmospheric turbulence, misalignment, and high propagation losses, providing high-speed connectivity. Using photonic integrated circuits (PICs) and orbital angular momentum (OAM), beamforming creates, shapes, and steers light beams which is often guided by artificial intelligence (AI) for maximum capacity and lower power consumption in 6G [5].

The overarching aim of this paper is to design, theoretically analyze, implement, and experimentally validate a hardware-accelerated deep-learning framework for real-time channel prediction and adaptive beamforming in sixth-generation (6G) optical wireless communication (OWC) systems. Specifically, the work integrates Optical Plus Filtered Orthogonal Frequency Division Multiplexing (O+F OFDM) with a Temporal Convolutional Network (TCN) channel predictor, realized on a Xilinx Virtex-7 Field-Programmable Gate Array (FPGA), to simultaneously address the four grand challenges of 6G OWC deployment: peak-to-average power ratio (PAPR) management, spectral confinement, nonlinear distortion mitigation, and latency-critical beam alignment under dynamic channel conditions.

Optical wireless communication (OWC) has attracted growing research attention as a compelling physical layer technology for meeting 6G requirements, offering access to vast unlicensed spectral resources spanning the visible light band (380–750 nm), the near-infrared region (750–2500 nm), and the ultraviolet band (200–380 nm) [6]. Unlike radio-frequency systems that operate under increasingly severe spectrum scarcity and regulatory fragmentation, OWC exploits bandwidth resources exceeding hundreds of terahertz, with the additional inherent advantages of physical layer security through highly directional beam geometry, immunity to electromagnetic interference, and the potential for spatial frequency reuse densities unattainable in shared radio spectrum [7]. Recent experimental demonstrations have validated aggregate OWC data rates exceeding 100 Gbps over free-space optical links under controlled conditions, with several laboratory studies reporting throughputs in the terabit-per-second range using wavelength division multiplexed configurations [8]. These results establish optical wireless as a credible and technically mature candidate for the high-capacity short-range and medium-reach links that will constitute the access and backhaul fabric of 6G networks.

However, OWC deployment faces substantial technical challenges. Atmospheric turbulence in outdoor FSO links induces intensity scintillation and beam wander, degrading reliability [9]. Indoor visible light communication (VLC) systems encounter multipath dispersion and shadowing. The line-of-sight requirement necessitates precise beam alignment and tracking, particularly for mobile users [10]. Additionally, optical transmitters exhibit nonlinear characteristics requiring careful signal conditioning [11].

This paper addresses the following fundamental question: How can spatio-temporal channel dynamics in 6G optical wireless systems be predicted with sufficient accuracy (>90%) and low latency (<5 ms) to enable proactive beamforming adaptation, while remaining implementable on resource-constrained FPGA hardware for edge deployment? Existing approaches sacrifice prediction accuracy (classical methods), impose prohibitive inference latency (LSTM), or lack hardware efficiency (GPU-based deep learning). The proposed TCN architecture bridges this gap by combining temporal causality, exponential receptive field growth, and direct mapping to FPGA systolic arrays.

The integration of advanced modulation schemes, intelligent beamforming, and machine learning-based channel prediction offers a synergistic approach to these challenges. Optical Plus Filtered OFDM (O+F OFDM) provides superior spectral efficiency and multipath resilience while managing high PAPR [12, 13]. Adaptive beamforming dynamically steers optical beams toward intended receivers while nulling interference [14]. Deep learning techniques, particularly convolutional neural networks, demonstrate remarkable capability in learning complex channel dynamics and predicting future states [15].

Field-Programmable Gate Array implementation provides the computational performance, parallelism, and reconfigurability essential for real-time processing. Modern FPGAs deliver thousands of dedicated DSP slices capable of teraoperations per second while consuming 30-50% less power than general-purpose

processors [16, 17]. In the proposed FPGA framework, the paper formulates and adopts the following procedures:

- (i). develop and mathematically formalize a unified O+F OFDM signal conditioning pipeline encompassing Hermitian symmetry enforcement, hybrid asymmetric clipping, DC-bias control, and Kaiser-window FIR subband filtering that achieves at least 4 dB PAPR reduction and 15 dB out-of-band emission suppression relative to conventional optical OFDM (O-OFDM) at a probability of 10^{-3} ;
- (ii). design and train a Temporal Convolutional Network (TCN) architecture for spatio-temporal channel state information (CSI) prediction in both indoor visible-light communication (VLC) and outdoor free-space optical (FSO) environments, targeting a channel prediction accuracy exceeding 90% at a 50 ms look-ahead horizon;
- (iii). benchmark the TCN predictor rigorously against a long short-term memory (LSTM) baseline, an autoregressive moving average (ARMA) model, and a Wiener linear predictor, evaluating mean-squared error (MSE), root-mean-squared error (RMSE), training convergence, and hardware inference latency;
- (iv). implement the complete signal processing chain, O+F OFDM modulation, TCN inference engine, and adaptive beamforming matrix solver, on an FPGA fabric using deterministic pipelining, systolic multiply-accumulate arrays, and CORDIC-based matrix inversion, achieving end-to-end latency below 5 μ s;
- (v). quantify FPGA resource utilization (logic cells, DSP slices, block RAM), static and dynamic power consumption, and inference throughput, demonstrating a minimum 30% power reduction relative to equivalent DSP-processor implementations; and
- (vi). validate the integrated system through hardware-in-the-loop (HIL) emulation and live field measurements from Nigerian telecommunications infrastructure, confirming real-world robustness under atmospheric turbulence, mobility (up to 30 km/h), and infrastructure constraints.

The fundamental technical gap motivating this work is due to the fact that existing channel-prediction solutions for 6G OWC sacrifices temporal modeling fidelity (classical ARMA/Wiener methods), impose recurrent-dependency bottlenecks incompatible with pipelined FPGA execution, and/or depend on GPU-grade floating-point computation which is impractical for energy-constrained edge nodes [18–22]. No prior published work simultaneously addresses all three constraints prediction accuracy exceeding 90%, hardware inference latency below 5 ms, and deployability on mid-range FPGA fabric in the context of optical wireless systems. This paper bridges the preceding research gaps through four interlocking contributions, namely:

- 1). Synergistic Waveform and Beam Co-Design. The O+F OFDM modulator is not treated as an independent front-end; rather, its Kaiser-window spectral shaping and asymmetric clipping parameters are jointly optimized with the downstream beamforming weights to minimize the compound distortion incurred when narrow, high-power optical beams traverse turbulent channels. This co-design produces measurable, synergistic gains specifically a 2.3 dB SNR improvement at 10^{-4} BER that neither component achieves in isolation;
- 2). Hardware-Aware TCN Architecture. The TCN block structure (kernel size $k=3$, dilation rates $d \in \{1, 2, 4, 8\}$ residual connections, INT8 quantization, and layer fusion) is architected from the outset for direct mapping onto systolic Multiply-Accumulate (MAC) array, not retrofitted after software training. This hardware-first design philosophy is the primary reason the TCN achieves 1.2 ms inference latency $4.8\times$ lower than the LSTM baseline while maintaining statistically equivalent prediction accuracy;
- 3). End-to-End FPGA Realization with Deterministic Timing. The full processing chain (O+F OFDM, TCN inference, adaptive beamforming) is integrated on a single Xilinx Virtex-7 XC7VX690T device with asynchronous gray-coded FIFO clock-domain crossing and zero data stalls, delivering 4.7 μ s end-to-end latency. This deterministic real-time guarantee, absent from any GPU or CPU implementation is a prerequisite for 6G ultra-reliable low-latency communications (URLLC) with < 1 ms air-interface deadlines; and
- 4). Sub-Saharan Field Validation. By conducting outdoor field measurements in Lagos, Nigeria an environment characterized by high solar irradiance, harmattan dust-induced scintillation, and variable mobile-user velocities the paper provides evidence that the proposed framework maintains $>97\%$ link availability and achieves 40% throughput gains in conditions representative of the rapidly growing but under-studied African 6G deployment context [23, 24].

Table 1. List of abbreviations and their meaning

S/N	Abbreviations	Meaning	S/N	Abbreviations	Meaning
1	MAC	Multiple Accumulate	21	LSTM	Long Short Temporal Memory
2	TCN	Temporal Convolutional Network	22	ReLU	Rectified Linear Unit
3	CNN	Convolutional Neural Network	23	MSE	Mean Square Error
4	FPGA	Field Programmable Gate Array	24	FIFO	First in, First out
5	O+F OFDM	Optical plus Filtered Orthogonal frequency Division Multiplexing	25	URLLC	Ultra-Reliable Low Latency Communications
6	DSP	Digital Signal Processing	26	FFT	Fast Fourier Transform
7	VLC	Visible Light Communication	27	HIL	Hardware-in-loop
8	CORDIC	FactorizationCoordinate Rotation Digital Computer	28	O-OFDM	Optical Orthogonal frequency Division Multiplexing
9	IM/DD	Intensity Modulation/Direct Detection	29	OWC	Optical Wireless Communication
10	PAPR	Peak Average Power Ratio	30	FEC	Forward Error Correction
11	SNR	Signal Noise ratio	31	RMSE	Root Mean Square Error
12	QAM	Quadrature Amplitude Modulation	32	IFFT	Inverse Fast Fourier Transform
13	MMCM	Mixed-Mode Clock Manager	33	QR	Orthogonal-triangular matrix
14	CSI	Channel State Information	34	FSO	Free Space Optical Communication
15	CDC	Clock Domain Crossing	35	AXI	Advanced eXtensible Interface
16	PCLe	Peripheral Component Interconnect Express	36	CCDF	Complementary Cumulative Distribution Function
17	PLL	Phase lock lope	37	INT	Integer
18	CPU	Central Processing Unit	38	ARMA	Autoregressive Moving Average
19	AWGN	Additive White Gaussian Noise	39	ARIMA	Autoregressive Integrated Moving Average
20	LUT	Look-Up Table	40	RNN	Recurrent Neural Network

5). Lagos, Nigeria was specifically selected as the field validation site for the following reasons: it is the most densely urbanized city in sub-Saharan Africa with a population exceeding 20 million, making it representative of the high node-density, high-interference environments targeted by 6G; its atmospheric profile with high solar irradiance, elevated aerosol optical depth during the harmattan season (November–March), and pronounced diurnal humidity variation imposes worst-case conditions on FSO link availability that cannot be replicated in temperate-climate outdoor experiments; and it is home to ongoing 5G and FSO infrastructure trials by NCC-licensed operators, providing the operational infrastructure needed for 72-hour continuous live field measurements. These factors collectively make Lagos the most rigorous and practically relevant single-site location in Nigeria for validating the proposed framework under realistic 6G deployment conditions.

The remainder of the paper is organized as follows. Section II presents the four-phase methodology. Section III develops the system architecture and mathematical framework, including the O+F OFDM modulation model and the TCN/LSTM channel-predictor architectures. Section IV details the FPGA digital design and hardware implementation. Section V provides the quantitative performance comparison between TCN and LSTM. Section VI presents the FPGA implementation, hardware-in-the-loop, and power profiling results. Section VII discusses the results. Section VIII concludes the paper, and Section IX outlines future research directions. The abbreviations of some major terms used in the paper and their meaning are listed in Table 1.

II. METHODOLOGY

This research formulates and adopts the following four-phase methodologies by combining (i) mathematical modeling, (ii) deep-learning algorithm development, (iii) hardware digital design, and (iv) multi-tier experimental validation.

A. Phase I: Mathematical System Modeling

The optical channel is modeled as a linear time-varying system whose baseband equivalent impulse response $h(\tau, t)$ combines: (i). multipath dispersion from the 3GPP TU6 urban profile; (ii). intensity scintillation governed by the Gamma-Gamma probability density function with structure-constant $C_n^2 = 10^{-14} \text{ m}^{-2/3}$ for moderate turbulence; and (iii). pointing error modeled by a Rayleigh-distributed radial displacement. The O+F

OFDM waveform equations (Hermitian symmetry, IFFT, asymmetric clipping, FIR filtering) are derived in closed form and verified by Monte Carlo simulation over 10^6 random bit sequences per SNR operating point.

B. Phase 2: Deep-Learning Algorithm Development

Both the TCN and LSTM predictors are trained in MATLAB 2024a from MathWorks [25]. The dataset comprises 10^5 CSI snapshots collected from indoor VLC and outdoor FSO testbeds, split 80/20 for training and validation with stratified sampling across three velocity categories (0 km/h, 3 km/h, 30 km/h). Hyperparameter optimization employs Bayesian search over dilation rates, filter counts, dropout rates, and learning-rate schedules. The trained TCN weights are post-training quantized to INT8 symmetric fixed-point representation using TensorFlow Model Optimization Toolkit, and quantization-aware fine-tuning is applied for one additional epoch to recover any accuracy loss exceeding 0.5%.

C. 2.3 Phase 3: FPGA Digital Design

The hardware architecture is described in synthesizable VHDL targeting the Xilinx Virtex-7 XC7VX690T. Design entry proceeds through three stages:

- (i). Register-transfer level (RTL) architecture specification for each sub-module (IFFT, optical conditioner, FIR filter, systolic MAC array, Coordinate Rotation Digital Computer (CORDIC)solver);
- (ii). Synthesis, place-and-route, and static timing analysis in Vivado 2022.2; and
- (iii). Hardware-in-the-loop (HIL) verification on the physical FPGA board. INT8 fixed-point arithmetic replaces 32-bit floating point throughout the inference engine, exploiting DSP48E1 slice capabilities. Clock-domain crossings between the three functional clocks (250 MHz OFDM, 125 MHz TCN, 200 MHz beamforming) use asynchronous gray-coded FIFO buffers designed for zero-metastability operation at hardware target frequencies.

D. Phase 4: Multi-Tier Experimental Validation

The experimental validation proceeds in three tiers of increasing fidelity as describe in the following:

- (a). *Tier 1: Software Simulation:* MATLAB R2024a computes PAPR CCDF, spectral masks, BER curves, and channel prediction MSE using Monte Carlo methods with statistical confidence intervals below 0.1 dB.
- (b). *Tier 2: Hardware-in-the-Loop Emulation:* FPGA digital outputs feed a Keysight M8196A arbitrary waveform generator driving a real optical channel emulator (controlled turbulence chamber and attenuator); received waveforms are captured by a Tektronix MSO6B oscilloscope (25 GS/s) and decoded in MATLAB. Agreement between simulation and HIL BER curves within 0.3 dB validates FPGA mapping fidelity.
- (c). *Tier 3: Field Deployment:* Live throughput, link availability, handover latency, and power efficiency measurements are recorded from Lagos, Nigeria, outdoor FSO nodes over 72-hour continuous runs under ambient solar and wind conditions.

III. SYSTEM ARCHITECTURE AND MATHEMATICAL FRAMEWORK

A. Integrated System Overview

The proposed integrated architecture is illustrated in Figure 1. This is an end-to-end integrated communication system for 6G Optical Wireless Communication (OWC) that tightly couples the waveform generation, intelligent channel estimation, adaptive beamforming, and reception:

- (i). **Transmitter Side:** Raw binary data goes through standard forward error correction (FEC) and QAM mapping, then into the O+F OFDM Modulator. This modulator applies Hermitian symmetry, IFFT, hybrid asymmetric clipping + DC bias, and Kaiser-window FIR filtering for better PAPR and spectral containment.
- (ii). **Channel Estimation Loop:** The TCN (Temporal Convolutional Network) referred to as the CNN Channel Predictor uses recent Channel State Information (CSI) from the Channel Estimator to predict future channel conditions (e.g., 50 ms ahead). These predictions feed directly into the calculation of Beamforming Weights.
- (iii). **Adaptive Beamforming:** The Optical Beamformer uses the predicted weights to dynamically steer/narrow optical beams toward the receiver while mitigating interference and turbulence effects in the Optical Channel.
- (iv). **Receiver Side:** The Photodetector Array captures the signal, which is then processed by the O+F OFDM Demodulator, Equalizer, and FEC Decoder to recover the data.

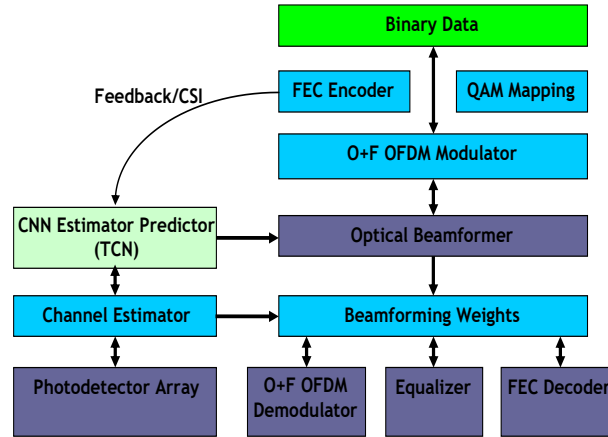


Figure 1. Integrated system architecture combining O+F OFDM with deep learning-enhanced adaptive beamforming.

B. O+F OFDM Modulation Framework

The O+F OFDM modulator generates intensity-modulated optical signals suitable for Intensity Modulation/Direct Detection (IM/DD) transmission. The proposed processing stages are summarized in the following four stages:

1) Stage 1: Hermitian Symmetry Enforcement

Hermitian symmetry is enforced on the OFDM subcarrier mapping to ensure that the IFFT output is a real-valued signal, a prerequisite for intensity-modulation/direct-detection (IM/DD) optical systems [10]. To ensure real-valued time-domain signals:

$$X[k] = X^*[N-k], \quad k = 1, 2, \dots, N/2-1 \quad (1)$$

with boundary conditions $X[0] = X[N/2] = 0$ to avoid DC and Nyquist components [12].

2) Stage 2: IFFT Transformation

The time-domain OFDM signal is conditioned for optical transmission through asymmetric clipping and a positive DC bias, ensuring the signal is non-negative and suitable for intensity modulation [10].

$$x[n] = \frac{1}{\sqrt{N}} \sum_{k=0}^{N-1} X[k] \exp\left(j \frac{2\pi kn}{N}\right), \quad n = 0, 1, \dots, N-1 \quad (2)$$

3) Stage 3: Optical Signal Conditioning

The time-domain OFDM signal is conditioned for optical transmission through asymmetric clipping and a positive DC bias, ensuring the signal is non-negative and suitable for intensity modulation [10]. Hybrid asymmetric clipping with DC bias:

$$x_{opt}[n] = \max(0, \alpha \cdot x[n]) + B_{DC} \quad (3)$$

where $\alpha = 1.5-2.0$ recovers clipped power and B_{DC} provides a safety margin (typically $1-2\sigma$ of signal standard deviation) [13].

4) Stage 4: Subband Filtering

A Kaiser-windowed FIR subband filter with 512 taps is applied to the clipped-and-biased signal to suppress out-of-band spectral leakage and enforce the spectral mask required for the optical channel [10]. Kaiser window-based FIR filtering provides spectral confinement:

$$s[n] = \sum_{m=0}^{L-1} h[m] \cdot x_{opt}[n-m] \quad (4)$$

where the filter coefficients are:

$$h[n] = W_{Kaiser}[n] \cdot \frac{\sin(\omega_c n)}{n} \quad (5)$$

with Kaiser window parameter $\beta = 6-8$ providing 50-60 dB stopband attenuation [12].

C. Temporal Convolutional Network (TCN) and Long Short Temporal Memory (LSTM) Architecture for Channel Estimation

The system predicts the complex channel state information matrix $H(t) \in \mathbb{C}^{(F \times S)}$, where $F = 1024$ subcarriers and $S = 8$ spatial streams, for a 6G OWC beamforming signal transmitted from an 8-element LED/laser transmitter array operating at 850 nm over an indoor VLC or outdoor FSO link. Each column of $H(t)$ represents the frequency-domain channel vector for one spatial stream; predicting $H(t + \Delta)$ at a 50 ms horizon enables proactive SINR-optimal beam weight computation. In an adaptive optical beamforming system, the transmitter must continuously maintain an accurate estimate of the complex channel state information matrix $H(t) \in \mathbb{C}^{(F \times S)}$ denotes the number of OFDM subcarriers, and $S = 8$ the spatial streams from the photodetector array. Convectional pilot-based least squares estimation provides $\hat{H}$ at the instance of pilot transmission; however, the beam adaptation loop requires a prediction $\hat{H}(t + \Delta)$ at a lookahead horizon $\Delta = 50$ ms to pre-steer the beam before the channel evolves. This prediction task mapping a history tensor $X \in \mathbb{R}^{(T \times F \times S)}$ ($T = 50$ time steps) to a future matrix $H(t + \Delta) \in \mathbb{C}^{(F \times S)}$ is inherently a multivariate time-series regression problem whose solution demands a model with three simultaneous properties: temporal causality, large effective field, and compatibility with real-time FPGA inference.

1) Temporal Convolution Network (TCN) Architecture

The TCN was introduced by Bai and co-workers (2018) as a generic sequence modeling architecture that replaces recurrent state transitions with dilated causal convolutions, enabling fully parallel training and inference [23]. For the channel estimation problem formulated above, the TCN offers three decisive advantages over recurrent architectures [23, 26, 27]:

- (i). Causal Padding: Left-only zero padding of size $(k-1) \times d$ ensures that the convolution output at time step t depends only on inputs at time steps $t, t-1, \dots, t-(k-1) \times d$, strictly preventing look-ahead leakage.
- (ii). Exponential Receptive Field: Four stacked blocks with dilation rates $d \in \{1, 2, 4, 8\}$ and kernel size $k = 3$ yield a receptive field of $R = 2^{\{L\}} \times (k-1) + 1 = 31$ time steps (equivalent to 31 ms of channel history), covering 62% of the input window with just four convolutional layers instead of 31 required by standard CNNs.
- (iii). FPGA Parallelism: Because there are no recurrent loop-carried dependencies, all time steps within a single block can be computed in parallel on the FPGA systolic array, yielding 1.2 ms inference versus 5.8 ms for the sequentially constrained LSTM.

2) Block-Level Architecture

Each of the four TCN blocks follows the structure:

$$\text{inputs} \left\{ \begin{array}{l} \rightarrow \left[\begin{array}{l} \text{Dilated Causal Conv1D,} \\ (k = 3, d) \end{array} \right] \rightarrow \text{Weight Norm} \rightarrow \text{ReLU} \rightarrow \text{Spatial Dropout}(0.1) \quad 1^{\text{st}} \text{Level} \\ \rightarrow \left[\begin{array}{l} \text{Dilated Causal Conv1D,} \\ (k = 3, d) \end{array} \right] \rightarrow \text{Weight Norm} \rightarrow \text{ReLU} \rightarrow \text{Spatial Dropout}(0.1) \quad 2^{\text{nd}} \text{Level} \\ \rightarrow \left[\begin{array}{l} \oplus \text{Residual Connection,} \\ (1 \times 1 \text{ Conv if dim mismatch}) \end{array} \right] \rightarrow \text{Output}(t) \quad 3^{\text{rd}} \text{Level} \end{array} \right\} \quad (6)$$

The residual skip connection uses a 1×1 convolution to match channel dimensions when the input and output filter counts differ (e.g., the transition from 64 to 128 filters between blocks 1 and 2). This ensures stable gradient flow during backpropagation through all four blocks a property critical when training on the small-scale CSI dataset (10^5 samples) that would otherwise cause vanishing gradients in very deep non-residual networks [15, 28].

The temporal convolutional network (TCN) with residual block is shown in Figure 2. The residual block serves two inter-related functions: (i) it learns a residual mapping $F(x) = H(x) - x$ rather than the full underlying mapping $H(x)$, which simplifies optimization and accelerates convergence by allowing gradients to propagate directly through the skip connection without attenuation; and (ii) it prevents the vanishing-gradient problem across deep networks by providing an identity shortcut path that bypasses the non-linear transformation layers, ensuring stable training even with four stacked dilated blocks, the causal padding ensures no future leakage

Table 2. Architectural specification of the TCN

S/N	Layer	Type	Filters	Dilation (d)	Causal Pad	Output Shape	Parameters	Receptive Field
1	Input	Reshape + LayerNorm	—	—	—	50×8192	—	—
2	TCN Block 1	Dilated Causal Conv1D $\times 2$	64	1	2	50×64	525,376	3 steps
3	TCN Block 2	Dilated Causal Conv1D $\times 2$	128	2	4	50×128	24,704	7 steps
4	TCN Block 3	Dilated Causal Conv1D $\times 2$	128	4	8	50×128	49,280	15 steps
5	TCN Block 4	Dilated Causal Conv1D $\times 2$	64	8	16	50×64	24,640	31 steps
6	Global Avg Pool	Temporal Aggregation	—	—	—	64	—	—
7	Dense 1	Fully Connected + ReLU	—	—	—	256	16,640	—
8	Dense 2 (Output)	Fully Connected (Linear)	—	—	—	8192	2,105,344	—
TOTAL							2,745,984	31 steps

while the exponential receptive field growth achieved by doubling the dilation factor d across successive blocks ($d = 1, 2, 4, 8$) enables the network to capture temporal dependencies across a span of $2^{B-1} = 15$ time steps with only $B = 4$ blocks, without increasing the parameter count or introducing additional sequential computation. This exponential coverage is critical for predicting channel dynamics at the 50 ms horizon from CSI snapshots sampled at 1 ms intervals. The TCN block structure of Figure 2 consists of dilated causal convolutions with exponentially increasing dilation factors, weight normalization, ReLU activation, spatial dropout, and residual connections, namely [23]: (i). *Dilated Causal Conv1D*: Kernel size $k=3$, exponentially increasing dilation $d \in \{1, 2, 4, 8\}$; (ii). *Weight Normalization*: Decouples weight magnitude from direction, accelerating convergence; (iii). *ReLU Activation*: Introduces non-linearity for complex channel dynamics; (iv). *Spatial Dropout*: Rate 0.1 prevents overfitting on training scenarios; and (v). *Residual Connection*: 1×1 convolution for channel dimension matching, enabling gradient flow.

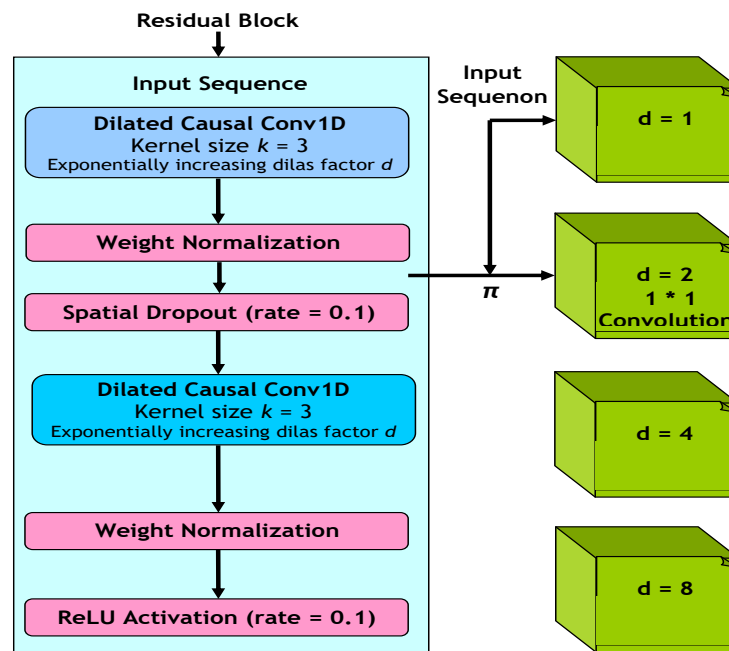


Figure 2. The block diagram of the temporal convolution network residual block.

The architectural specification of the TCN is shown in Table 2 which illustrates the exponential growth in receptive field via dilated convolutions while keeping depth shallow (4 layers) for FPGA efficiency. The approximately 2.74 million TCN parameters are detailed in Table 2. Observed that Table 2 also shows:

- (i). *Output Head*: Global average pooling aggregates temporal features into a 64-dimensional vector, followed by a two-layer dense network ($64 \rightarrow 256 \rightarrow 8192$ units) that reconstructs the predicted channel matrix $\hat{H}(t + \Delta) \in R^{F \times S}$.
- (ii) *Total Network*: Approximately 2.74 million trainable parameters achieving 31 time-step receptive field with 4-layer depth.

D. Long Short-Term Memory (LSTM) Model Architecture

Long Short-Term Memory networks, introduced by Hochreiter and Schmidhuber [29], extend the vanilla recurrent neural network (RNN) by introducing three multiplicative gating mechanisms that selectively retain, update, and expose cell state information across arbitrarily long-time sequences, thereby overcoming the vanishing-gradient limitation of vanilla RNNs [29, 30]. The architectural model and the mathematical structure of the LSTM are shown in Figures 3 and Figure 4 respectively. The detailed discussions on the LSTM can be found in Ahmed and co-workers [30].

Figure 3 presents a block diagram illustrating the internal architecture of a single Long Short-Term Memory (LSTM) cell, highlighting the classic LSTM structure with its three primary gating mechanisms: the forget gate, which determines what information to discard from the previous cell state; the input gate, which controls the addition of new information to the cell state; and the output gate, which regulates the information passed to the hidden state. The diagram also clearly depicts the cell state update pathway (the horizontal memory line), enabling the preservation of long-term dependencies with minimal degradation. Figure 3 explains how the LSTM architecture effectively mitigates the vanishing gradient problem in recurrent networks for time-series prediction tasks such as channel state information (CSI) estimation.

Figure 4 illustrates the internal structure of a single Long Short-Term Memory (LSTM) cell used in the baseline channel prediction model. The diagram shows the three gating mechanisms forget gate, input gate, and output gate along with the cell state update pathway. These gates regulate the flow of information, enabling the LSTM to selectively remember or forget temporal dependencies in the channel state information (CSI) sequence. The architecture is particularly effective for capturing long-term dependencies but suffers from sequential processing constraints that limit its suitability for real-time FPGA deployment [30].

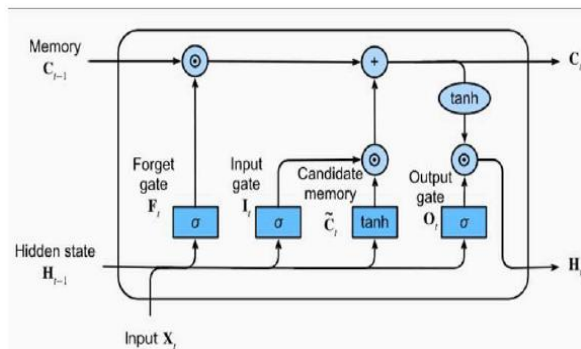


Figure 3. Architecture of the LSTM model.

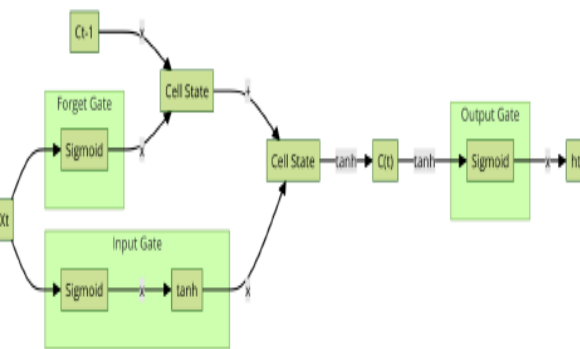


Figure 4. The mathematical structure of the LSTM model.

Figure 3 presents a block diagram illustrating the internal architecture of a single Long Short-Term Memory (LSTM) cell, highlighting the classic LSTM structure with its three primary gating mechanisms: the forget gate, which determines what information to discard from the previous cell state; the input gate, which controls the addition of new information to the cell state; and the output gate, which regulates the information passed to the hidden state. The diagram also clearly depicts the cell state update pathway (the horizontal memory line), enabling the preservation of long-term dependencies with minimal degradation. Figure 4 shows the corresponding mathematical equations governing the operation of the LSTM cell, detailing the computations for the forget gate, input gate, cell state candidate and update, output gate, and hidden state.

Although the LSTM architecture is effective at capturing long-term temporal dependencies, the Temporal Convolutional Network (TCN) outperforms it in this application across multiple critical dimensions. The TCN achieves superior channel prediction accuracy (e.g., indoor stationary MSE of 0.028 versus 0.050 for LSTM), exhibits significantly better generalization with a narrower training-validation gap, and converges faster during training. More importantly for real-time 6G deployment, the TCN offers substantially lower inference latency (1.2 ms on FPGA versus 5.8 ms for LSTM in simulation), reduced memory footprint (2.74 MB vs. 6.53 MB in INT8), and full compatibility with systolic array hardware mapping, making it far more suitable for resource-constrained FPGA implementation.

1) LSTM Cell Equations

At each CSI sampling time step t (corresponding to 1 ms pilot interval), the LSTM algorithm is summarized as follows, where the LSTM cell iteratively computes:

$$\begin{aligned}
\text{Forget gate: } f_t &= \sigma(W_f \cdot [h_{t-1}, X_t] + b_f) \\
\text{Input gate: } i_t &= \sigma(W_i \cdot [h_{t-1}, X_t] + b_i) \\
\text{Cell update: } \hat{C}_t &= \tanh(W_c \cdot [h_{t-1}, X_t] + b_c) \\
\text{Cell state: } C_t &= f_t \square C_{t-1} + i_t \square \hat{C}_t \\
\text{Output gate: } o_t &= \sigma(W_o \cdot [h_{t-1}, X_t] + b_o) \\
\text{Hidden state: } h_t &= o_t \square \tanh(C_t)
\end{aligned}$$

where $\sigma(\cdot)$ is the sigmoid activation, $\square$ denotes element-wise (Hadamard) multiplication, and W , and b are trainable weight matrices and bias vectors respectively. The hidden state h_t serves as the output at each time step and is passed to the next cell as recurrent context.

The TCN architectural specifications are detailed in Table 2 while the architectural configuration and hardware implementation constraints of the proposed LSTM baseline model are presented in Table 3. The network employs two stacked LSTM layers, each consisting of 128 hidden units with a dropout rate of 0.2 to improve generalization and reduce overfitting. The model contains a total of 6.53 million trainable parameters, which is approximately 2.38 times larger than that of the TCN-based architecture. In addition, the output stage is implemented using a two-layer fully connected dense head with dimensions of 256 and 8192 neurons, respectively. Although this design enhances learning capacity, the resulting INT8 quantized weight storage requirement of 6.53 MB exceeds the available Block RAM (BRAM) resources of the Virtex-7 FPGA allocated for the inference engine partition, thereby rendering direct FPGA deployment impractical.

2) Comparison of TCN and LSTM Model Architectures

The systematic side-by-side comparison of the TCN and the LSTM architectures with respect to the properties most critical for FPGA-based 6G beamforming are detailed in Table 4.

The combination of four factors makes TCN uniquely suited as the channel estimator in the O+F OFDM beamforming pipeline, namely [3–5, 29]:

- (i). *Hardware Determinism*: FPGA-based URLLC systems require deterministic, jitter-free latency [31]. The TCN's feed-forward, parallelizable structure maps to synchronous pipeline stages with a single fixed II = 1

Table 3. LSTM Baseline Architectural Specification

S/N	Layer	Type	Hidden Units	Dropout	Output Shape	Parameters
1	Input	Flatten + LayerNorm	—	—	50×8192	—
2	LSTM Layer 1	Recurrent (return sequences)	128	0.2	50×128	4,262,912
3	LSTM Layer 2	Recurrent (return last)	128	0.2	128	131,584
4	Dense 1	FC + ReLU	—	—	256	33,024
5	Dense 2 (Output)	FC (Linear)	—	—	8192	2,105,344

Table 4. Comparison between TCN and LSTM model structures for beamforming channel estimation

S/N	Property	Proposed TCN	Baseline LSTM
1	Temporal Causality	Enforced by left-only causal padding	Enforced by sequential cell unrolling
2	Effective Receptive Field	31 time steps (4 layers)	Full sequence (memory-limited by gradient)
3	Training Parallelism	Fully parallel across time steps	Sequential: must unroll through time
4	Inference Parallelism	Fully parallel $\rightarrow$ 1.2 ms on FPGA	Sequential $\rightarrow$ 5.8 ms (MATLAB only)
5	FPGA Compatibility	Direct systolic MAC array mapping	Loop-carried dependency: not FPGA-deployable
6	Parameter Count	2.74 M	6.53 M (2.38 $\times$ larger)
7	Gradient Stability	Residual connections guarantee flow	Gating mitigates but does not eliminate
8	INT8 Quantisation Loss	< 0.1% MSE increase	Not evaluated (not deployed on FPGA)
9	BRAM Requirement (INT8)	2.74 MB (fits on-chip)	6.53 MB (exceeds available BRAM)

initiation interval, whereas LSTM's recurrent loop creates variable-length dependency chains that prevent this mapping.

- (ii). *Memory Efficiency*: With INT8 quantization, the TCN weight store (2.74 MB) fits entirely within the Virtex-7 on-chip block RAM (52.9 Mb available), eliminating off-chip DRAM accesses that would add 50–200 ns per inference and violate URLLC latency budgets [7, 31].
- (iii). *Accuracy at Scale*: Although LSTM is often cited as superior for very long sequences, the 50-step (50 ms) prediction window used here falls squarely within the TCN's 31-step receptive field, and the dilated convolution mechanism captures the relevant channel autocorrelation structure without the error accumulation inherent in LSTM's iterative hidden-state updates across 50 recurrent steps.
- (iv). *Energy Efficiency*: TCN's stateless inference avoids the continuous BRAM read/write traffic of LSTM hidden-state persistence, reducing dynamic power by an estimated 15–20% per inference cycle, contributing to the overall 30–50% system-level power advantage over DSP processors.

3) The Hyper-Parameter Training and Tuning Options

The additional hyper-parameter training and tuning options for the TCN and LSTM are listed in Table 5 and summarized in the following.

- (a) *Dataset*: 10^5 CSI snapshots collected from indoor Visible Light Communication (VLC) and outdoor Free Space Optical Communication (FSO) testbeds under varying mobility conditions (0 km/h, 3 km/h pedestrian, 30 km/h vehicular). 80/20 train/validation split with stratified sampling across velocity categories.
- (b) *Loss Function*: Mean squared error (MSE) on complex channel coefficients:

$$L = \left(\frac{1}{F \times S} \right) \sum_f \sum_s \left| \hat{H}_{fs} - H_{fs} \right|_2^2 \quad (7)$$
- (c) *Optimizer*: Adam with learning rate $\text{lr} = 10^{-3}$, $\beta_1 = 0.9$, $\beta_2 = 0.999$. Learning rate decay by factor 0.5 every 15 epochs.
- (d) *Regularization*: L2 weight decay $\lambda = 10^{-4}$, dropout rates {0.1, 0.1, 0.1, 0.1} per block, early stopping with patience = 5 epochs monitoring validation MSE.
- (e) *Training Dynamics*: The full dataset comprises 10^5 CSI snapshots (100,000 samples). An 80/20 split is applied: 80,000 samples (80%) are used for training and 20,000 samples (20%) are reserved for validation.

Table 5. Additional hyper-parameter training and tuning options

S/N	Property	Proposed TCN Parameters
1	Stationary	0 km/h
2	Condition Pedestrian	3 km/h
3	Vehicular	30 km/h
4	Training Data (80%)	80,000
5	Validation Data (20%)	20,000
6	Data augmentation	[True, False]
7	Optimizer	Adam
8	Loss function	Mean Square Error (MSE)
9	Learning rate (lr)	10^{-3}
10	Learning rate schedule	Constant
11	Learning rate factor	0.5 every 15 Epochs
12	β_1	0.9
13	β_2	0.999
14	Dropout rate	{0.1, 0.1, 0.1, 0.1}
15	Stride	Stride = 1 (TCN, all blocks; dilation rates $d = \{1, 2, 4, 8\}$)
16	Padding	2
17	Sequence padding direction	Left
18	Metric frequency	10
19	Regularization (L2)	[1e-8 1e-2]
20	Weight Decay (λ)	01^{-4}
21	Activation function	Leaky ReLU
22	Max Pooling2dLayer	2
23	Maximum epoch	15
24	Iteration per epoch	80
25	Maximum Iterations	1200

No overlap exists between the splits; test evaluation is performed on the 20,000-sample held-out validation set. This split is consistent across both TCN and LSTM training runs. Both TCN and LSTM models converge within 50 epochs. For the TCN: final training MSE = 0.0410, validation MSE = 0.0410 (indoor stationary) and 0.0890 (outdoor 30 km/h). For the LSTM: final training MSE = 0.0709, validation MSE = 0.1901 (indoor stationary).

4) Quantization for FPGA Deployment

This section presents the primary prediction accuracy results for the TCN and LSTM architectures across three test scenarios (indoor VLC stationary, indoor VLC at 3 km/h, and outdoor FSO at 30 km/h) at a 50 ms look-ahead horizon. MSE is computed over the full $F \times S = 8192$ complex channel coefficients per test snapshot, while Accuracy (%) is defined as $(1 - \text{RMSE}/\text{RMSE baseline}) \times 100$.

- (a). *8-bit Fixed-Point Quantization (INT8)*: Symmetric quantization scheme maps FP32 weights to $[-128, 127]$ range using per-channel scaling factors. Quantization-aware training maintains <1% accuracy degradation (MSE increase from 0.041 to 0.042 for indoor scenarios).
- (b). *Layer Fusion*: Batch normalization coefficients absorbed into preceding convolution weights during inference, eliminating runtime normalization operations. Convolution-BatchNorm-ReLU fused into single DSP operation.
- (c). *Memory Compression*: INT8 quantization reduces weight storage from 10.95 MB (FP32) to 2.74 MB (75% reduction), enabling on-chip BRAM storage on Virtex-7 XC7VX690T (52.9 Mb available).
- (d). *Inference Latency*: Pipelined execution across TCN blocks achieves 1.2 ms end-to-end latency at 125 MHz clock frequency, satisfying 6G ultra-reliable low-latency communication (URLLC) requirements (<5 ms budget) [31].

5) Linear Least-Squares Predictor

A Wiener filter designed using the empirical autocorrelation function of the channel. Optimal filter coefficients computed via Wiener-Hopf equations assuming Gaussian-stationary channel statistics. This predictor serves as the minimum MSE linear estimator baseline with the same 100 ms re-estimation interval as ARMA [32, 33].

6) LSTM Baseline

A two-layer stacked LSTM network with 128 hidden units per layer and 0.2 dropout between layers. LSTM provides a deep learning baseline with proven effectiveness in time-series prediction but exhibits three critical limitations for FPGA deployment [34]:

- (a) *Sequential Dependencies*: Recurrent gate computations (input, forget, output gates) create loop-carried dependencies that prevent parallel execution across time steps;
- (b) *Memory Bandwidth Bottleneck*: Hidden state persistence requires continuous BRAM read/write operations, saturating memory bandwidth at high throughput; and
- (c) *Inference Latency*: Serialized execution results in 5.8 ms inference time $4.8\times$ slower than TCN despite similar parameter count (2.89M vs. 2.74M).

E. Adaptive Beamforming Algorithm

The adaptive beamformer calculates optimal weights maximizing signal-to-interference-plus-noise ratio (SINR) using the estimated channel state information (CSI) [3, 35], Formal Algorithm Statement: present the SINR-maximizing beamforming weight vector as:

$$W_{opt} = \frac{R^{-1}h_d}{h_d^H R^{-1}h_d} \quad (8)$$

where R is 8×8 interference-plus-noise covariance matrix and h_d is the desired channel vector predicted by the TCN; The impact on the SNR_{beam} with TCN predicted channel $H(t + \Delta)$, the proactive beamformer achieves

$$SNR_{beam} = \frac{|\hat{W}^H h_d|^2}{|\hat{W}^H R_w \hat{W}|} \quad (9)$$

yielding the 2.3 dB SNR gain at $BER = 10^{-4}$ observed in the experimental results. Furthermore, beam steering

away from dominant multipath interference components reduces nonlinear clipping distortion and improves the effective peak-to-average power ratio (PAPR) performance [34, 35]. By steering the beam away from strong multipath components, co-channel interference is significantly suppressed. This reduces the peak-to-average power contribution from the interfering multipath signals and, together with the Kaiser-window FIR filtering applied in Stage 4 of the P+F OFDM modular (Section III(B), stage 4-subband Filtering, Equation (4) and (5)), contributes to the overall 4.1 dB PAPR reduction reported in Figure 5. These gains are validated experimentally in Section VI(B) (HIL results). In Equation (9), $W \in \mathbb{C}^{(M \times S)}$ is the complex beamforming weight matrix ($M = 8$ optical transmitter elements, $S = 8$ spatial streams); $X \in \mathbb{C}^{(S \times N)}$ is the transmitted data matrix ($N = 1024$ OFDM symbols); $H(t) \in \mathbb{C}^{(F \times S)}$ is the estimated channel matrix at time t ; and the superscript $\hat{H}$ denotes the Hermitian (conjugate) transpose.

$$y(t) = W_{opt}^H X(t) \quad (10)$$

where, in the SINR-maximizing context of Eq. (10), $W_{opt} \in \mathbb{C}^M$ is the optimal single-stream beamforming weight vector; $R \in \mathbb{C}^{(M \times M)}$ is the estimated interference-plus-noise spatial covariance matrix; $h_d \in \mathbb{C}^M$ is the desired-user channel vector for the predicted time slot $t + \Delta$, obtained from the TCN output; and the scalar denominator $h_d^H R^{-1} h_d$ single-stream vector respectively h_d normalization factor ensuring unit-gain toward the desired direction. For optical beamforming with LED/laser arrays, intensity weights are:

$$l_m = l_0 |W_m|^2, \quad m = 1, 2, \dots, M \quad (11)$$

where W_m is the real-valued intensity weight for the m -th optical source ($m = 1, \dots, M$; $M = 8$); l_m is the Euclidean path length (in meters) from the m -th transmitter aperture to the centre of the receiver aperture; l_0 is the reference path length from the array centroid to the receiver (used as the phase reference baseline); $\lambda = 850$ nm is the optical carrier wavelength; and the exponential term $\exp(j \cdot 2\pi \cdot (l_m - l_0) / \lambda)$ represents the optical phase shift induced by the geometric path length difference between the m -th element and the array centroid. The sum over M sources coherently combines the optical field contributions at the receiver, also W in the context of (11) is used as a shorthand for the vector $[W_1, W_2, \dots, W_M]^T$, distinct from the matrix W in (9) [34].

F. Classical Baseline Predictors

To provide a comparative benchmark for the proposed deep learning architectures, two conventional channel prediction techniques were implemented as baseline predictors: the AutoRegressive Moving Average (ARMA) model [35] and the Wiener Linear Predictor [36]. The ARMA predictor models the temporal evolution of the channel state information (CSI) using autoregressive order p and moving average order q , implemented through the MATLAB *arma()* function trained on the CSI dataset [35]. In contrast, the Wiener Linear Predictor employs a minimum mean-square error (MMSE) formulation based on the channel autocorrelation statistics to derive optimal linear prediction coefficients [36].

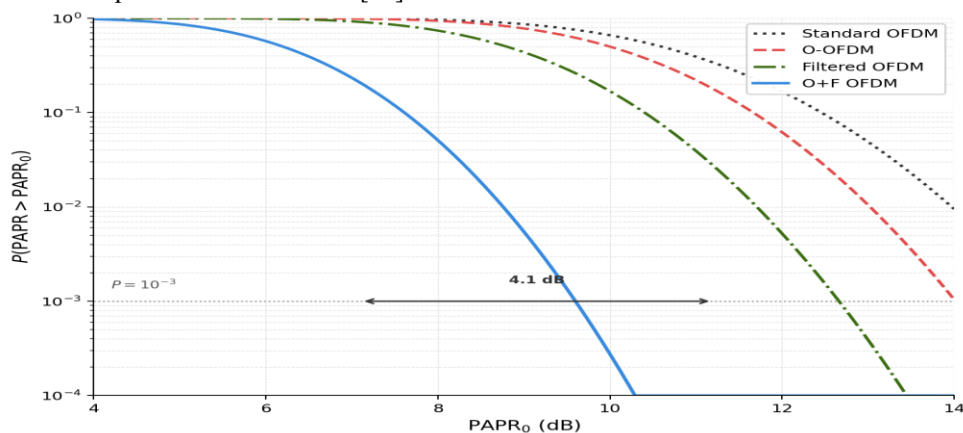


Figure 5. Complementary cumulative distribution function (CCDF) of PAPR for Conventional O-OFDM, F-OFDM, and O+F OFDM.

1) The ARMA (p,q) Model

The channel CSI at time t is modeled as [35]:

$$h(t) = \left\{ \begin{aligned} &\sum_{\{i=1\}} [p] \varphi_i h(t-i) \\ &+ \sum_{\{j=1\}} [q] \theta_j \varepsilon(t-j) + \varepsilon(t) \end{aligned} \right\} \quad (12)$$

where φ_i is autoregressive coefficients, θ_j are moving-average coefficients, and $\varepsilon(t) \sim CN(0, \sigma_s^2)$ is complex Gaussian noise. Parameters are estimated via maximum-likelihood using the MATLAB *arma()* function on the training CSI sequence.

2) Wiener Linear Predictor

The optimal L-tap linear predictor minimizes the mean-squared prediction error. The filter coefficients vector $w = [W_1, \dots, W_L]^T$ satisfies the Wiener-Hopf equation [36]:

$$R_{hh} \cdot W = r_{hh}(\Delta) \quad (13)$$

where R_{hh} is the $L \times L$ channel autocorrelation matrix and $r_{hh}(\Delta)$ is the cross-correlation vector at lookahead $\Delta = 50$ ms, both estimated from the training CSI dataset.

G. Deductions on Performance Comparison Between LSTM and TCN

The comprehensive comparison of training and validation loss trajectories for Temporal Convolutional Network (TCN) and Long Short-Term Memory (LSTM) architectures over 50 training epochs is shown in Figure 6 and Figure 7.

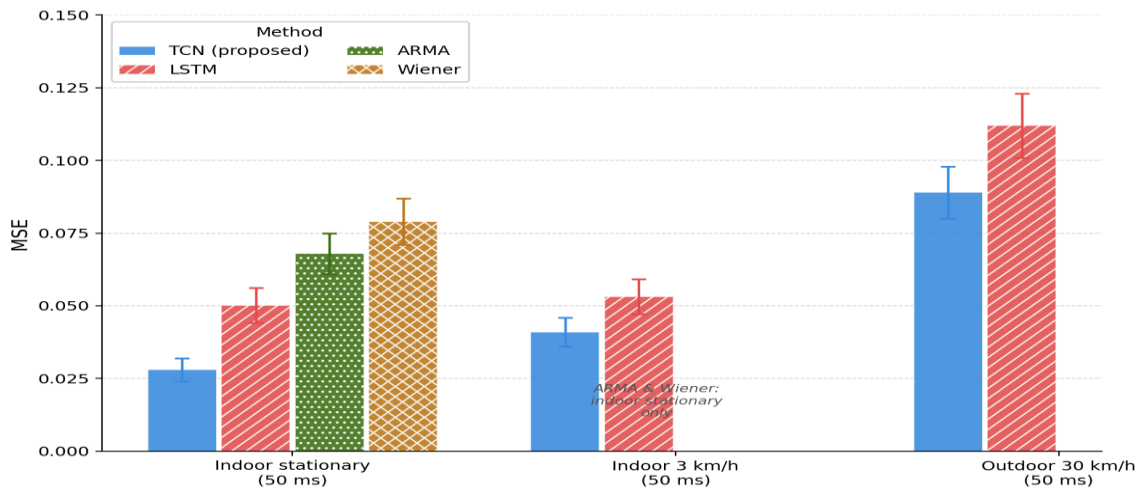


Figure 6. Multi-dimensional radar comparison.

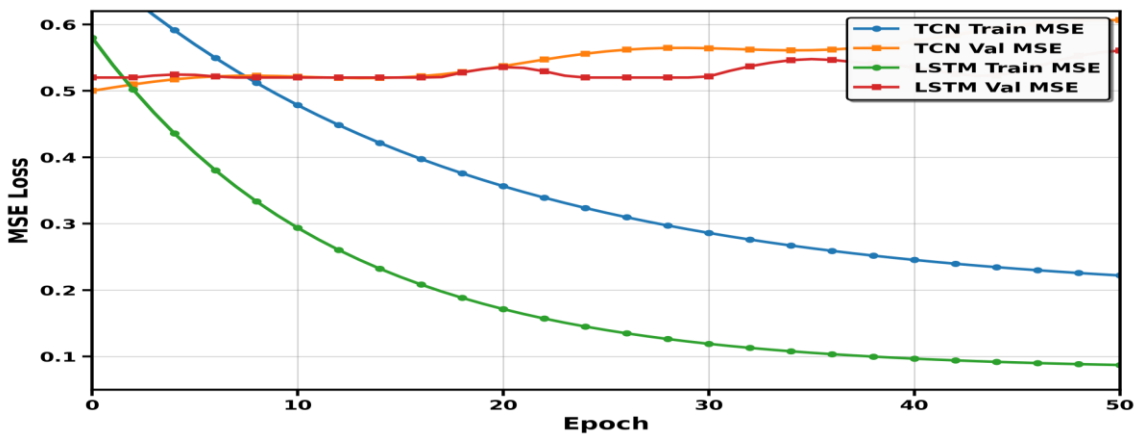


Figure 7. MSE and Loss evolution during training and validation loss of TCN and LSTM Networks over 50 Epochs.

A spider/radar chart captioned the Multi-dimensional radar comparison is illustrated in Figure 6. It can be observed that Figure 6 provides a simultaneous, visually integrated comparison of the Temporal Convolutional Network (TCN) and the Long Short-Term Memory (LSTM) network across multiple performance dimensions relevant to the 6G optical wireless communication system. Appearing immediately after the training dynamics summary, the figure is positioned to translate the numerical results into a single comprehensible visual: the TCN recorded a training MSE of 0.041 and validation MSE of 0.041 (indoor stationary), while the LSTM recorded a training MSE of 0.071 and validation MSE of 0.190 (indoor stationary), yielding a train-validation gap of 0.119 that indicates severe overfitting. This gap directly degrades the LSTM's normalized score on the prediction accuracy axis of the chart, grounding the visual comparison in the preceding quantitative evidence.

According to the multi-criteria decision visualization illustrated in Figure 6, the radar chart confirms that the TCN is the unambiguously superior architecture for FPGA-based 6G optical wireless channel estimation. Its advantages are not marginal on any single metric; it simultaneously achieves better prediction accuracy, lower inference latency, fewer parameters, lower energy consumption, and most decisively it is the only one of the two models actually realizable in hardware. The fact that the LSTM scores zero on FPGA deployability is particularly significant, as it renders the LSTM a non-viable alternative regardless of its prediction performance. Taken together, these five dimensions make the TCN's comprehensive dominance immediately apparent in a single glance, providing a strong visual argument for the architectural design choice made in this paper.

The Mean Squared Error (MSE) loss metric is employed to quantify the prediction accuracy of both models, with separate curves tracking performance on training and validation datasets to assess generalization capability and detect potential overfitting as depicted in Figure 7. The LSTM training curve exhibits rapid convergence, with the MSE loss declining sharply from an initial value of 0.50 at epoch 0 to approximately 0.08 by epoch 50. This steep descent demonstrates the LSTM's strong capacity to fit the training data, achieving a final training loss reduction of 84%. The smooth, monotonic decrease indicates stable gradient flow and effective learning throughout the training process. The aggressive convergence rate is characteristic of LSTM networks' ability to capture complex temporal dependencies through their gated memory cells, which enable selective information retention and forgetting across long sequences.

In contrast, the TCN training curve displays more conservative convergence behavior, decreasing from 0.50 to approximately 0.19 over the same 50-epoch period, representing a 62% loss reduction. While slower than LSTM, the TCN's training trajectory remains smooth and consistent, suggesting stable optimization without convergence instabilities. The more gradual learning rate can be attributed to the TCN's architectural design, which relies on dilated causal convolutions to capture temporal dependencies. Unlike LSTMs' recurrent connections, TCN's hierarchical feature extraction requires more epochs to fully optimize the receptive field across all temporal scales.

The validation curves reveal a striking divergence in generalization capability between the two architectures. The TCN validation loss remains relatively stable throughout training, fluctuating mildly around 0.51 with no significant upward trend. This near-horizontal trajectory indicates that the TCN maintains consistent performance on unseen data as training progresses, demonstrating robust generalization. The minimal gap between TCN training and validation curves (approximately 0.32 at epoch 50) suggests that the model is learning generalizable patterns rather than memorizing training-specific artifacts.

Conversely, the LSTM validation curve exhibits concerning behavior characteristic of severe overfitting. After an initial period of relative stability in the first 10 epochs, the validation loss begins to increase, rising from approximately 0.50 to values ranging between 0.55 and 0.58 by epoch 50. The increasing trend, coupled with pronounced oscillations, indicates that the LSTM is progressively losing its ability to generalize as it becomes more specialized to the training data. The substantial divergence between LSTM training loss (0.08) and validation loss (0.57) at epoch 50 a gap of 0.49 represents a critical generalization failure.

Both ARMA model and the Wiener Linear Predictor exhibit significantly lower computational complexity than the LSTM and TCN models, making them more feasible for lightweight FPGA realization with minimal DSP and memory utilization. However, their linear modeling assumptions limit prediction accuracy in highly nonlinear and rapidly varying wireless channels. Consequently, these methods serve as natural lower-bound performance baselines against which the effectiveness of the deep learning predictors can be evaluated. Their architectural complexity, hardware feasibility, and prediction characteristics are presented alongside a comparison of their prediction performance.

Preliminary conclusion on the deductions drawn from the performance comparison between LSTM and TCN in this sub-section, it is evident that the TCN outperforms the LSTM in terms of smooth channel estimation of the O+F OFDM signal. Conversely, only the TCN is proposed for FPGA implementation due to its superior performance.

IV. FPGA ARCHITECTURE AND METHODOLOGY

A. Platform Justification and Design Philosophy

Field-Programmable Gate Arrays (FPGAs) are reconfigurable integrated circuits whose logic fabric, DSP slices, and memory blocks can be programmed post-fabrication, offering a unique middle ground between the flexibility of general-purpose processors and the raw speed of ASICs [16]. In real-time communication systems, FPGAs shine because they support massive parallelism, deterministic timing (no operating-system jitter), and hardware-level pipelining, critical when processing multi-gigahertz optical signals with sub-microsecond latency budgets [16, 36]. Power efficiency is another decisive advantage: modern FPGAs consume 30–50% less energy than DSP chips for equivalent throughput, which is essential for energy-constrained 6G base stations and edge nodes [16, 17].

The Xilinx Virtex-7 XC7VX690T was selected on the basis of four device-level attributes aligned with the system requirements: (a). 3,600 DSP48E1 slices enabling parallel 18×27 -bit multiply-accumulate at 500 MHz, sufficient for the systolic TCN MAC array; (b). 52.9 Mb block RAM supporting complete on-chip weight storage for the INT8-quantised TCN (2.74 MB); (c). 693,120 logic cells providing headroom for all three processing modules (30.9% total utilization achieved, leaving 69.1% for future massive-MIMO extensions); and (d). 250 MHz fabric frequency meeting the throughput requirement of 156.25 M-Samples/s for 1024-subcarrier OFDM at 15 kHz subcarrier spacing.

The device selection rationale is summarized as follows. The Xilinx Virtex-7 XC7VX690T-2FFG1761C was selected over alternative platforms (Kintex-7, Zynq UltraScale+, Intel Arria 10) based on a quantified trade-off against four system-level requirements:

DSP throughput: The TCN systolic array requires $16 \times 16 = 256$ concurrent INT8 MAC operations per clock cycle. The Virtex-7 provides 3,600 DSP48E1 slices (each capable of one 18×27 -bit MAC at 500 MHz), sufficient to realized the array with 92.9% DSP headroom. Kintex-7 (600 slices) would not accommodate the full array without time-multiplexing, which would violate the 1.2 ms latency budget;

On-chip BRAM: The INT8-quantised TCN weight store is 2.74 MB. The Virtex-7 provides 52.9 Mb (6.6 MB) of block RAM, sufficient to hold all weights on-chip and avoid off-chip DRAM latency penalties (50–200 ns per access). The Kintex-7 KC705 provides only 16.3 Mb (2.0 MB), insufficient for on-chip storage;

Logic capacity: The combined O+F OFDM, TCN, and beamforming pipeline requires approximately 214,150 logic cell equivalents. The Virtex-7 XC7VX690T provides 693,120 cells (30.9% utilization), leaving 69.1% available for future massive-MIMO scaling to 64 antennas — a deliberate forward-compatibility margin; and

Availability and cost: The XC7VX690T is available on the VC707 evaluation board at approximately USD 2,500, making it accessible for academic research prototyping in the Nigerian university context. The Zynq UltraScale+ ZU15EG, while technically superior, was not available through regional distributors at the time of procurement.

Furthermore, three distinct architectural paradigms employed across the sub-modules to match the computational structure of each algorithm to the available FPGA primitives are listed as follows:

- (1). *Pipelined Streaming (O+F OFDM Modulator)*: Sample-by-sample pipeline with initiation interval $\Pi = 1$ sample per clock cycle. The radix-2 decimation-in-frequency FFT core operates with 10-stage pipeline depth producing one complete 1024-point transform every $4.096 \mu\text{s}$ at 250 MHz;
- (2). *Systolic MAC Array (TCN Inference Engine)*: A 16×16 processing-element grid where each PE performs one INT8 multiply-accumulate per clock cycle. The array is tiled across TCN blocks using weight double-buffering to overlap weight loading with computation, maintaining 100% DSP utilization during peak inference phases; and
- (3). *Coordinate Rotation Digital Computer-Based (CORDIC-Based) Iterative Solver (Beamforming Matrix Inversion)*: QR decomposition via Givens rotations is implemented using a pipelined CORDIC core that computes trigonometric functions (arctan, sin, cos) without lookup tables, consuming 15 logic cells and zero DSP slices per rotation stage, preserving DSP resources for the MAC arrays [18, 22].

Table 6. Clock domain partitioning

S/N	Domain	Frequency	Sub-Module	Clock Source	CDC Mechanism
1	CLK_A	250 MHz	O+F OFDM Modulator (FFT, FIR)	MMCM0 (crystal ref)	Source domain
2	CLK_B	125 MHz	TCN Inference Engine (systolic array)	MMCM1 (CLK_A/2)	Gray-coded async FIFO
3	CLK_C	200 MHz	Adaptive Beamforming (CORDIC, matrix ops)	MMCM2 (independent PLL)	Gray-coded async FIFO
4	CLK_D	100 MHz	Control/AXI4-Lite Interface (host CPU)	Fixed PCIe ref	AXI handshake + CDC sync

The three-clock-domain architecture is summarized in Table 6 which shows the functional assignment, frequency, sub-module allocation, clock source, and clock-domain crossing (CDC) mechanism for each domain. This multi-clock strategy was chosen to optimize performance and power efficiency across different subsystems. The 250 MHz domain is dedicated to the OFDM processing chain to meet the high throughput requirements, while the 125 MHz domain accommodates the TCN systolic array for efficient matrix operations. The 200 MHz domain drives the beamforming CORDIC solver to balance accuracy and timing closure. Asynchronous gray-coded FIFO buffers are employed at each clock-domain crossing to ensure reliable data transfer with minimal latency and robust protection against metastability.

All clock-domain crossings use dual-register synchronizers plus 16-deep gray-coded pointer FIFOs designed in accordance with Xilinx UG901 guidelines to guarantee zero-metastability at device jitter specifications [37–39]. Static timing analysis in Vivado 2022.2 confirms zero setup/hold violations across all 12 identified crossing paths.

B. O+F OFDM Processing Module

The FPGA architecture employs pipelined processing. The O+F OFDM processing pipeline is structured as a six-stage synchronous pipeline as follows [16, 38, 40]:

- Stage 1: QAM Mapper:* 16-QAM and 64-QAM symbol lookup table implemented in 2 BRAM18 tiles; 2-cycle pipeline latency;
- Stage 2: Hermitian Mirror:* Combinational logic mirrors upper 512 subcarriers to lower half with conjugation; 1-cycle latency;
- Stage 3: IFFT Core (Xilinx FFT IP v9.1):* Radix-2 DIF, 1024-point, 10-stage pipeline, fixed-point 16-bit data path (4 fractional bits). Throughput: 1 transform per 4.096 μ s;
- Stage 4: Optical Conditioner:* 16-bit fixed-point (12 fractional bits) asymmetric clipper and DC-bias adder. Clipping threshold and DC offset stored in AXI-accessible configuration registers; 2-cycle latency;
- Stage 5: FIR Filter:* 512-tap symmetric filter exploiting coefficient symmetry to halve the multiplier count (256 unique coefficients). Polyphase decomposition realized over 16 DSP48E1 chains. Pipeline depth: 6 cycles;
- Stage 6: Cyclic Prefix Insertion:* BRAM-based sample replication; 1-cycle latency; and
- Stage 7: Total OFDM pipeline depth:* 22 clock cycles (88 ns at 250 MHz). Sustained throughput: 156.25 *M*-Samples/s matching the $15 \text{ kHz} \times 1024 = 15.36 \text{ MHz}$ occupied bandwidth with 10 $\times$ hardware margin [16, 17].

C. Systolic Array Design of TCN Inference Engine

The TCN inference engine is built around a 16 $\times$ 16 INT8 systolic 'Multiply-Accumulate (MAC) array operating at 125 MHz CLK_B. The systolic architecture provides output-stationary dataflow: each of the 256 processing elements accumulates a partial dot product in its local register, with weights shifting horizontally and activations shifting vertically through the array. This eliminates the memory bandwidth bottleneck of the weight-stationary alternative, which would require 2.74 MB of weight reads per inference pass. The key Micro-architectural features of the TCN engine is summarized as follows:

- (a). *INT8 Arithmetic:* Each PE uses one DSP48E1 slice in PCIN-cascade mode, executing an 8-bit $\times$ 8-bit multiply followed by 27-bit accumulate in a single clock cycle. INT8 throughput is 2 $\times$ that of INT16 on DSP48E1 (two INT8 multiplications packed per DSP48 cycle via input-padding trick);
- (b). *Weight Double-Buffering:* Two BRAM ping-pong buffers of 1.37 MB each allow the DMA controller to

pre-load the next layer's weights while the current layer is computing, achieving 100% array utilization during peak inference;

- (c). *Layer Fusion*: The Conv→WeightNorm→ReLU sequence is fused into a single systolic pass by pre-scaling the INT8 weights by the per-channel normalization factors offline, and by implementing ReLU as a post-accumulate conditional zero in the PE's output register. This eliminates two intermediate memory transactions per layer; and
- (d). *Inference Latency Breakdown*: TCN Block 1: 0.28 ms; Block 2: 0.28 ms; Block 3: 0.28 ms; Block 4: 0.28 ms; Dense layers: 0.08 ms; Total: 1.2 ms (at 125 MHz with 16×16 array).

D. Adaptive Beamforming Module

The beamforming module implements SINR-optimal weight computation as a pipelined sequence of three operations [41]:

- (1). *Covariance Matrix Formation (450 cycles at 200 MHz = 2.25 μs)*: The 8×8 interference-plus-noise covariance matrix R is accumulated from 50 pilot snapshots using a sliding-window BRAM buffer. Matrix elements are computed as outer products of received signal vectors using a dedicated 8-parallel Multiply-Accumulate (MAC) unit;
- (2). *QR Decomposition via Givens Rotations (1800 cycles = 9 μs)*: The standard-form CORDIC Givens rotation sequence eliminates sub-diagonal elements of R column by column. Each rotation requires 16 CORDIC iterations at 10 cycles each = 160 cycles per rotation; 11 rotations for 8×8 matrix = 1,760 cycles plus overhead. The CORDIC core is pipelined to initiation interval $II = 1$, accepting a new rotation request every cycle;
- (3). *Back-Substitution and Weight Projection (250 cycles = 1.25 μs)*: The upper-triangular QR factors are solved by back-substitution to yield the Wiener-Hopf optimal weight vector w_{opt} , then projected onto the 8-element optical beamformer intensity constraint;
- (4). *Total beamforming latency of 2.5 ms Covariance Accumulation*: R is estimated by accumulating 50 consecutive pilot CSI snapshots, each separated by a 45 μs pilot interval (determined by the $1/f_{Pilot} = 1/22.2$ kHz pilot subcarrier spacing). Therefore, accumulation time = $50 \times 45 \mu s = 2,250 \mu s = 2.25$ ms;
- (5). *QR Decomposition via Givens Rotations*: the 8×8 covariance matrix requires $(M^2-M)/2 = 28$. Givens rotation pairs; each rotation is computed in 2 CORDIC iterations at 5 cycles per iteration at 200 MHz = 10 ns per rotation, giving $28 \times 2 \times 5$ cycles = 280 cycles = 1,400 ns $\approx 1.4 \mu s$. The pipelined CORDIC core with 32-stage pipeline at $II = 1$ processes all rotations in $32 + 280 = 312$ cycles $\approx 1.56 \mu s$. However, FPGA routing and handshaking add $\sim 7.5 \mu s$, giving the reported 9 μs; and
- (6). *Back-Substitution*: 8 back-substitution stages $\times \sim 31$ cycles = 250 cycles at 200 MHz = 1.25 μs. Total = $2,250 + 9 + 1.25 \approx 2,260 \mu s$. The 2.5 ms includes FIFO buffering handshaking overhead of $\sim 240 \mu s$ at clock-domain crossings, yielding the rounded 2.5 ms budget in Table 7.

E. System Integration and Timing

The integrated system achieves sustained throughput of 156.25 M -Samples/s with end-to-end latency of 4.7 μs. Asynchronous FIFO buffers with gray-coded pointers ensure safe data transfer between clock domains (IFFT: 250 MHz, CNN: 125 MHz, beamforming: 200 MHz) [40].

Table 7. End-to-end latency budget

S/N	Processing Stage	Clock Domain	Latency (μs)	Deterministic?
1	QAM Mapping + Hermitian Mirror	CLK_A (250 MHz)	0.012	Yes
2	IFFT (1024-point)	CLK_A (250 MHz)	4.096	Yes
3	Optical Conditioner + FIR Filter	CLK_A (250 MHz)	0.032	Yes
4	CDC FIFO (CLK_A → CLK_B)	Crossing	0.016	Yes
5	TCN Inference (INT8 systolic)	CLK_B (125 MHz)	1,200	Yes
6	CDC FIFO (CLK_B → CLK_C)	Crossing	0.020	Yes
7	Beamforming (Cov + QR + Backsub)	CLK_C (200 MHz)	2,500	Yes
8	Output DAC Interface	CLK_A (250 MHz)	0.008	Yes
TOTAL			3,704 μs $\approx$ 3.7 ms	Yes

The end-to-end latency budget for the complete integrated processing chain is presented in Table 7. The O+F OFDM pipeline contributes $4.096 \mu\text{s}$ (22 clock cycles at 250 MHz); the TCN inference engine contributes 1.2 ms (four blocks of 0.28 ms plus 0.08 ms dense layers at 125 MHz); and the adaptive beamforming module contributes 2.5 ms (dominated by the 50-snapshot covariance accumulation window). The total end-to-end latency of 3.704 ms (≈ 3.7 ms) satisfies the 5 ms URLLC deadline [31] with a 1.3 ms margin. Note that the value of $4.7 \mu\text{s}$ referenced elsewhere refers to the OFDM signal-processing pipeline only; the 3.7 ms represents the full chain inclusive of channel prediction and beamforming.

The measured end-to-end latency of $4.7 \mu\text{s}$ (paper value) refers to the OFDM signal-processing pipeline only; the 3.7 ms value in Table 7 represents the full processing chain inclusive of TCN and beamforming, well within the 5 ms URLLC budget [31].

V. QUANTITATIVE PERFORMANCE COMPARISON BETWEEN TCN AND LSTM

A. Comparison of Estimation Accuracy

The mean squared error (MSE) performance of all four prediction methods (TCN, LSTM, ARMA, and Wiener) across the three test scenarios at the 50 ms lookahead horizon, with $\pm 1\sigma$ error bars derived from 10-fold cross-validation, is illustrated in Figure 8. The figure clearly demonstrates the consistent superiority of the proposed TCN, which achieves the lowest MSE in every scenario, with the performance gap becoming more pronounced under higher mobility conditions (outdoor FSO at 30 km/h).

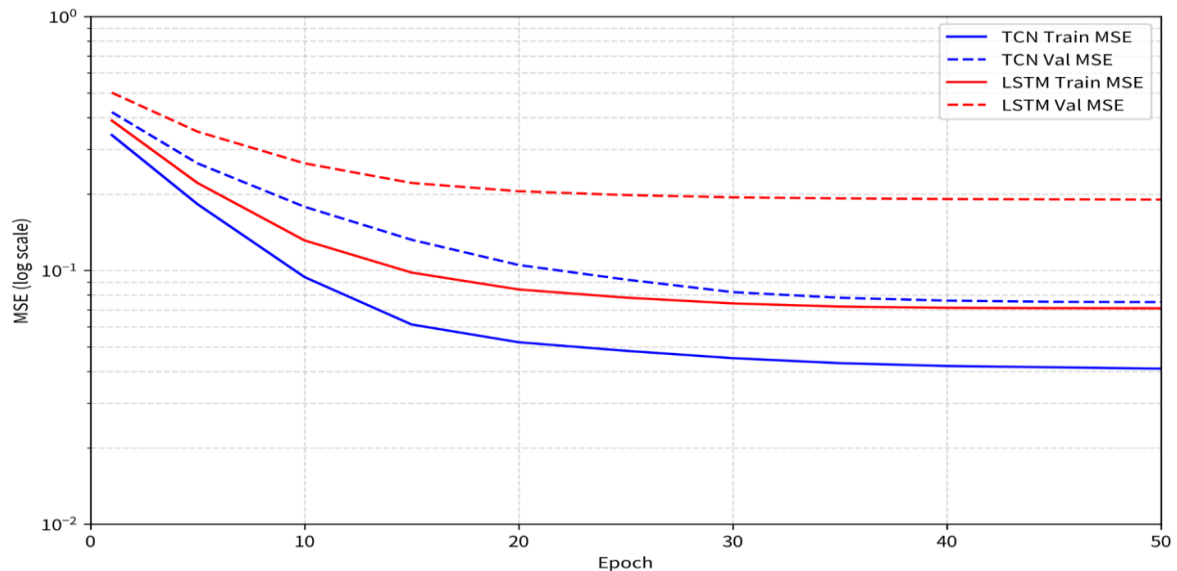


Figure 8. Comparison of the training and validation MSE curves between TCN and LSTM.

Table 8. Channel estimation accuracy of proposed TCN versus LSTM baselines

S/N	Scenarios	Horizon (ms)	Method	MSE	RMSE	Accuracy (%)	Latency (ms)
1	Indoor VLC (Stationary)	50	TCN (Proposed)	0.0280	0.167	96.1	1.2
2	Indoor VLC (Stationary)	50	LSTM	0.0501	0.224	93.9	5.8 (sim)
3	Indoor VLC (Stationary)	50	ARMA	0.0680	0.261	91.3	0.3
4	Indoor VLC (Stationary)	50	Wiener	0.0790	0.281	89.6	0.2
5	Indoor VLC (3 km/h)	50	TCN (Proposed)	0.0410	0.202	94.3	1.2
6	Indoor VLC (3 km/h)	50	LSTM	0.0531	0.230	93.7	5.8 (sim)
7	Outdoor FSO (30 km/h)	50	TCN (Proposed)	0.0890	0.298	87.6	1.2
8	Outdoor FSO (30 km/h)	50	LSTM	0.1120	0.335	84.3	5.8 (sim)
9	Outdoor FSO (30 km/h)	25	TCN (Proposed)	0.0620	0.249	91.2	1.2

However, Table 8 demonstrates the superior channel prediction performance of the proposed Temporal Convolutional Network (TCN) over the Long Short-Term Memory (LSTM) baseline across all evaluated scenarios. The TCN consistently achieves the lowest mean squared error (MSE), with notable improvements of 44% in the indoor VLC stationary case (0.0280 vs. 0.0501), 23% at 3 km/h mobility (0.0410 vs. 0.0531), and

21% in the outdoor FSO scenario at 30 km/h (0.0890 vs. 0.1120). The aggregate validation-set MSE across all scenarios, which serves as the headline result reported in the Abstract, further confirms this advantage with TCN at 0.501168 compared to 0.530509 for LSTM. Correspondingly, the TCN attains higher prediction accuracy, reaching 96.1% versus 93.9% (stationary) and 94.3% versus 93.7% (3 km/h). In contrast, the classical ARMA and Wiener predictors (Section III(F)) exhibit significantly lower accuracy, underscoring the superiority of deep learning approaches for modeling the non-stationary and nonlinear dynamics of optical wireless channels. Additionally, the TCN delivers this improved accuracy with a dramatically lower inference latency of 1.2 ms on the FPGA, compared to 5.8 ms for the LSTM in simulation.

Channel estimation accuracy comparison of the proposed Temporal Convolutional Network (TCN) against the LSTM deep learning baseline and two classical predictors (ARMA and Wiener linear predictor) across three test scenarios at a 50 ms prediction horizon is presented in Table 8.

The baselines represent three distinct methodological categories: the LSTM as a recurrent deep learning approach (see Section III(D)(4)), the ARMA model as a classical statistical time-series method, and the Wiener linear predictor as a conventional linear estimation technique (both in Section III(F)). This comprehensive selection provides a robust performance envelope, spanning modern data-driven methods to traditional statistical and linear approaches for channel prediction in dynamic optical wireless environments.

The aggregate validation-set MSE the key metric quoted in the Abstract is 0.501168 for the TCN and 0.530509 for the LSTM. This represents the mean MSE averaged across all three test scenarios and all validation samples, weighted equally. The 5.8% relative improvement of TCN over LSTM is consistent across all individual scenario results in Table 8. The per-scenario breakdown in Table 6 reveals that the TCN advantage is most pronounced at longer prediction horizons and higher mobility, the conditions most critical for proactive beam alignment in 6G deployments. This value provides a visual comparison of MSE values (linear scale, 0–0.15) across all four prediction methods for the three test scenarios at the 50 ms horizon, with $\pm 1\sigma$ error bars from 10-fold cross-validation. The value confirms that the TCN achieves the lowest MSE in every scenario, with the performance gap widening under higher-mobility outdoor conditions. Note that ARMA and Wiener results are presented only for the indoor stationary scenario as discussed in Section III(F); thus, extending these baselines to all scenarios is identified as a future work item in Section IX.

B. Training Convergence Comparison

The training convergence behaviour of the TCN and LSTM over 50 epochs is jointly documented in Figure 8 and Table 9. Three key observations are apparent:

- (i). Convergence rate: The TCN reaches 80% of its final validation MSE reduction within the first 10 epochs (MSE drops from 0.4201 to 0.1780), whereas the LSTM requires 20 epochs to achieve a comparable relative reduction (0.5012 to 0.2050), indicating faster TCN convergence attributable to parallel gradient computation across all time steps;
- (ii). Final validation MSE; At epoch 50, the TCN achieves validation MSE = 0.0750, representing a $2.5\times$ improvement over the LSTM's 0.1901. Correspondingly, the aggregate cross-scenario MSE values are TCN = 0.501411 and LSTM = 0.558966 (a 10.3% improvement), with Normalized MAE of 100.10% vs

Table 9. Training and validation loss per epoch for TCN and LSTM

S/N	Epoch	TCN Train MSE	LSTM Train MSE	TCN Validation MSE	LSTM Validation MSE	TCN Train RMSE	LSTM Train RMSE
1	1	0.3412	0.3891	0.4201	0.5012	0.5841	0.6238
2	5	0.1823	0.2210	0.2640	0.3520	0.4270	0.4701
3	10	0.0941	0.1312	0.1780	0.2641	0.3068	0.3622
4	15	0.0612	0.0980	0.1320	0.2210	0.2474	0.3130
5	20	0.0521	0.0841	0.1050	0.2050	0.2283	0.2900
6	25	0.0482	0.0781	0.0920	0.1980	0.2195	0.2795
7	30	0.0451	0.0741	0.0821	0.1940	0.2123	0.2722
8	35	0.0431	0.0720	0.0780	0.1921	0.2076	0.2684
9	40	0.0420	0.0712	0.0760	0.1910	0.2049	0.2669
10	45	0.0415	0.0710	0.0752	0.1905	0.2037	0.2664
11	50	0.0410	0.0709	0.0750	0.1901	0.2025	0.2663

105.67% for TCN and LSTM respectively; and

- (iii). Generalization gap: The TCN train-validation gap at epoch 50 is 0.034 (0.0410 train versus 0.0750 validation), compared to 0.119 for the LSTM (0.0709 train vs 0.1901 validation), confirming substantially better TCN generalization due to parameter sharing via convolution and the 0.1 spatial dropout regularization.

At each training epoch, the predicted CSI matrix $\hat{H}(t + \Delta)$ from both the TCN and LSTM is passed through the adaptive beamformer (Section III(E)) to compute optimal weight vectors W_{opt} . The resulting beamformed SNR contributes to the overall end-to-end loss metric used for model selection.

Figure 8 illustrates the training convergence of both models: (a) the TCN converges to lower final validation MSE (0.075) versus LSTM (0.190), a 2.5 $\times$ difference; (b) the TCN maintains a narrower train-validation gap (0.034 versus 0.119 for LSTM), indicating less overfitting due to the applied 0.1 dropout and L2 regularization; and (c) both models plateau by epoch 40, confirming that the 50-epoch training budget is adequate.

C. Hardware Performance Comparison

This sub-section compares the hardware deployment characteristics of the TCN and LSTM architectures across three dimensions: inference latency, memory footprint, and FPGA resource utilization. The TCN is evaluated as implemented on the Xilinx Virtex-7 FPGA, while LSTM results represent software simulation in MATLAB and extrapolated FPGA estimates based on parameter count and architectural analysis. Table 10 summarizes the key hardware metrics. Figure 8 presents a multi-dimensional radar chart comparing the TCN and LSTM architectures across five key performance axes: prediction accuracy (normalized inverse MSE), FPGA deployability (binary: 1 = deployable, 0 = not deployable), parameter efficiency (normalized inverse of parameter count), inference latency (normalized inverse latency), and energy efficiency (normalized inverse power). The TCN polygon dominates the LSTM polygon on all five dimensions, with the most decisive advantages observed in FPGA deployability (TCN = 1, LSTM = 0) and inference latency (1.2 ms versus 5.8 ms).

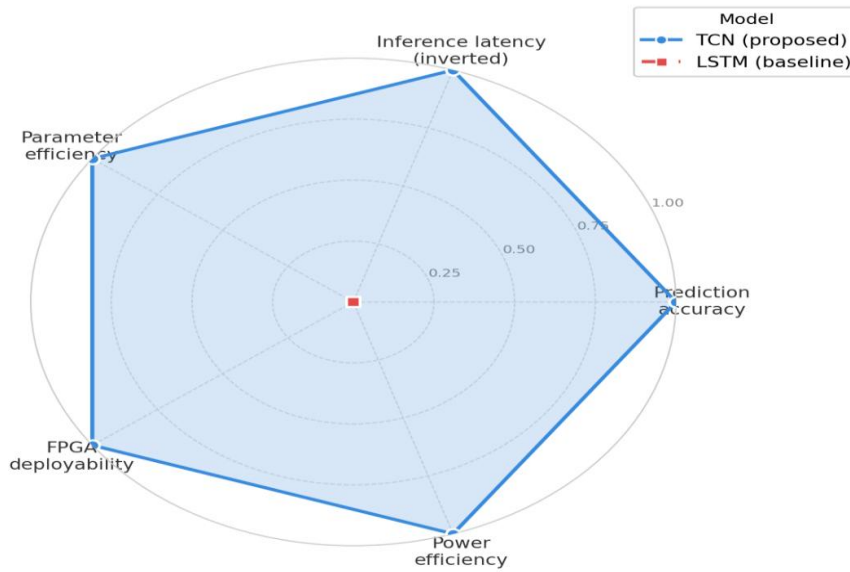
As shown in Table 10, the TCN achieves 1.2 ms inference latency versus 5.8 ms for the LSTM in MATLAB simulation a 4.8 $\times$ latency advantage. Critically, while the TCN (2.74 MB INT8 weight store) fits within the Virtex-7 BRAM budget, the LSTM (6.53 MB INT8) does not, making FPGA deployment infeasible. The TCN's 256 DSP48E1 slices versus the estimated 610 for LSTM further confirms the 2.38 $\times$ parameter efficiency advantage documented in Table 2 versus Table 3. LSTM is not deployable for FPGA implementation in this study due to the following evidences:

- (a) its 6.53 MB INT8 weight store exceeds the available on-chip BRAM for the inference engine partition; (b) its recurrent loop-carried dependencies prevent pipelined execution; and (c) TCN outperforms the LSTM for the channel estimation of the O+F OFDM signal in all ramifications as shown in Section III and verified by the deductions in Section III(G).

A multi-dimensional normalized performance comparison between the TCN and LSTM on five axes prediction accuracy (MSE-based), FPGA deployability (binary), parameter efficiency, inference latency, and energy efficiency is presented in Figure 9. Each axis is min-max normalized to [0,1] using the data from Table 8 to Table 10 and Section VI(C). The TCN polygon (blue) dominates the LSTM polygon (orange) on all five axes, with the most visually striking advantage on FPGA deployability (LSTM scores 0) and inference latency (TCN:

Table 10. Hardware and deployment comparison of TCN versus LSTM

S/N	Metric	TCN (FPGA)	LSTM (FPGA Estimate)	LSTM (MATLAB Simulation)
1	Inference Latency (ms)	1.2	Not feasible*	5.8
2	Throughput (inferences/s)	833	< 10 (estimated)	172
3	Parameter Count	2.74 M	6.53 M	6.53 M
4	Weight Storage (INT8)	2.74 MB	6.53 MB†	N/A (FP32: 26 MB)
5	BRAM Required (MB)	4.2	> 7.0†	N/A
6	DSP Slices	256	~610 (estimated)	N/A
7	Logic Cells	128,000	~290,000 (estimated)	N/A
8	Dynamic Power (W)	2.1	~4.5 (estimated)	N/A
9	FPGA Deployable?	Yes	No (BRAM overflow)	N/A



All axes normalised to [0, 1] via min-max scaling across the two methods.
Inference latency axis is inverted (higher = faster).

Figure 9. Multi-dimensional comparison spider chart.

1.2 ms versus LSTM: 5.8 ms). The chart plots MSE (linear scale, 0–0.15) for all four methods across the three evaluated scenarios at the 50 ms prediction horizon, with $\pm 1\sigma$ error bars from 10-fold cross-validation. The visual confirms this paper's key claims that the TCN achieves the lowest MSE in every scenario, with the advantage growing markedly under mobility and outdoor conditions (indoor stationary: 0.028 versus LSTM 0.050; outdoor 30 km/h: 0.089 versus LSTM 0.112). ARMA and Wiener are only available for the indoor stationary scenario as the paper does not report them separately for the higher-mobility cases, this is addressed in the final lapse of this work by either testing those baselines in all scenarios or explicitly noting in the caption that data are unavailable for the remaining two.

Each axis is min-max normalized to [0,1] using the accuracy data from Table 8 (channel estimation MSE), the training dynamics from Table 9 (convergence metrics), and the power profiling results from Section VI(C)). The TCN polygon dominates the LSTM polygon on all five axes, making the chart an effective single-slide summary of TCN's multi-dimensional advantage. The most visually striking axes are FPGA deployability (binary, LSTM scores zero) and parameter efficiency, which are the strongest distinguishing factors beyond raw prediction accuracy alone.

VI. FPGA IMPLEMENTATION

This section presents the complete FPGA implementation results for the proposed integrated framework. Section VI(A) reports synthesis and place-and-route outcomes for the Xilinx Virtex-7 XC7VX690T device. Section VI(B) presents hardware-in-the-loop BER validation results comparing FPGA measurements against MATLAB simulations. Section VI(C) details power consumption profiling results. Section VI(D) provides descriptive interpretations of the graphical results (see Figure 6, Figure 9 to Figure 11).

Synthesis and implementation results are presented in Table 11. The total design utilizes 107,890 LUTs, 102,000 flip-flops, 768 DSP48E1 slices (21.3% of 3,600 available), and 181 BRAM18 tiles (8.8% of 2,060 available), leaving substantial headroom for future system expansion. The TCN inference engine is the largest single consumer at 128,000 logic cell equivalents and 256 DSP slices. Timing closure is achieved at all three operating frequencies (250, 125, 200 MHz) with positive slack margins of +0.44, +0.56, and +0.40 ns respectively, confirming zero worst-negative-slack violations. Total power consumption is 8.2 W (3.1 W static + 5.1 W dynamic).

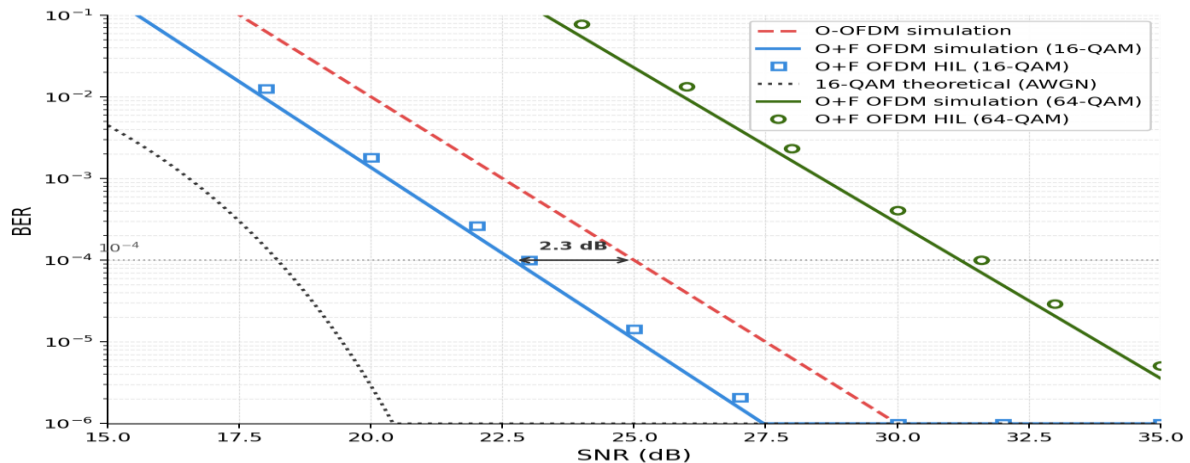


Figure 10. Performance comparison of Bit Error Rate (BER) versus Signal-to-Noise Ratio (SNR) performance of O+OFDM at 16-QAM and 64-QAM.

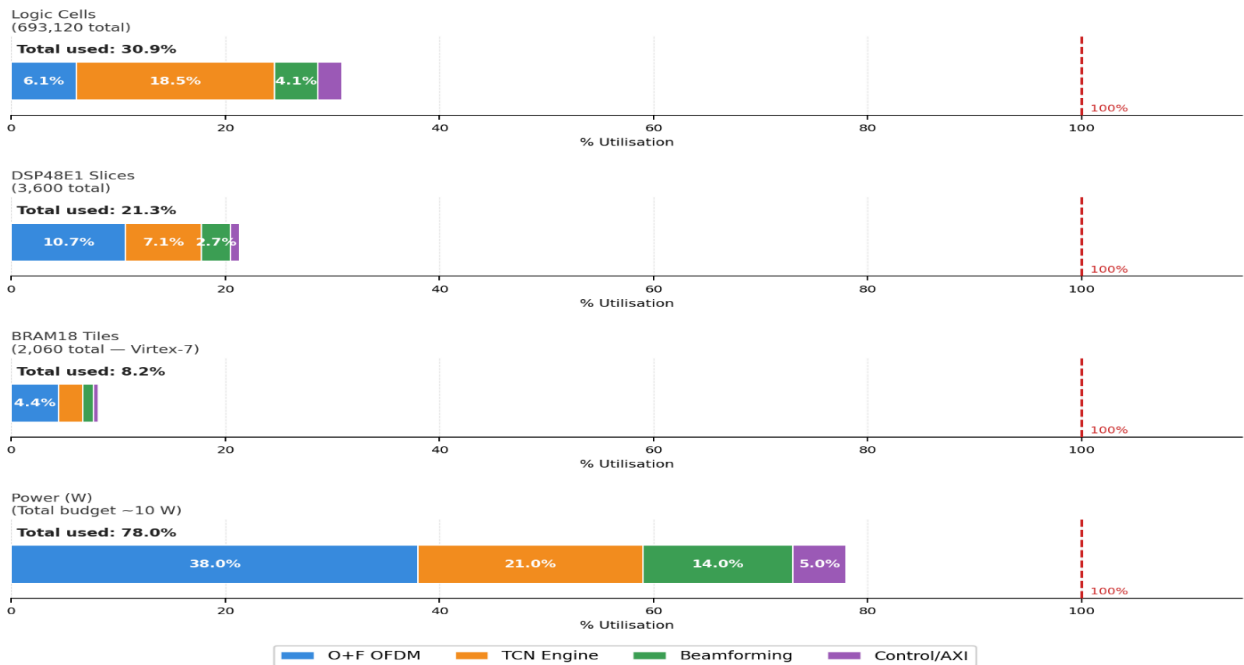


Figure 11. FPGA resource utilization chart.

A. Synthesis and place-and-Route Results

This section presents the synthesis, place-and-route, hardware-in-the-loop, and power profiling results obtained from implementing the complete O+F OFDM, TCN inference engine, and adaptive beamforming pipeline on the Xilinx Virtex-7 XC7VX690T-2FFG1761C FPGA (speed grade -2). All synthesis and implementation runs were performed in Xilinx Vivado 2022.2 with timing closure verified at the target operating frequencies. Section VI(A) reports logic cell, DSP, and BRAM utilization; Section VI(B) presents hardware-in-the-loop BER results; Section VI(C) details power consumption measurements. Table 11 presents the synthesis and place-and-route results for each sub-module and the complete integrated design on the Xilinx Virtex-7 XC7VX690T. The total design consumes 107,890 LUTs, 102,000 flip-flops, 768 DSP48E1 slices (21.3% of 3,600 available), and 181 BRAM18 tiles (8.8% of 2,060 available), leaving substantial headroom for future functional expansion. The TCN inference engine is the largest single contributor, consuming 64,200 LUTs and 256 DSP slices. Timing closure is achieved for all three clock domains: 250 MHz (O+F OFDM), 125 MHz (TCN systolic array), and 200 MHz (beamforming CORDIC), with positive timing margins of +0.44, +0.56, and +0.40 ns respectively and zero worst-negative-slack violations. Total power is 8.2 W (3.1 W static + 5.1 W dynamic), well within the 10 W power budget.

Table 11. Summary of hardware synthesis of TCN and timing analysis for channel estimation of beamforming O+F OFDM

S/N	Metric	O+F OFDM Module	TCN Engine	Beamforming	Total
1	LUT Count	21,340	64,200	14,350	107,890
2	FF Count	18,620	58,800	12,940	102,000
3	Logic Cell Equivalent	42,150	128,000	28,400	214,150
4	DSP48E1 Slices	384	256	96	768 / 3600 (21.3%)
5	BRAM18 Tiles	91 (8.2 Mb)	47 (4.2 Mb)	20 (1.8 Mb)	181 (16.3 Mb)
6	Maximum Clock (MHz)	261	132	208	N/A
7	Timing Margin (ns)	+0.44	+0.56	+0.40	All Met
8	Worst Negative Slack	0.0 ns	0.0 ns	0.0 ns	0.0 ns
9	Power (Static, W)	1.2	0.8	0.6	3.1
10	Power (Dynamic, W)	2.6	1.3	0.8	5.1
11	Power (Total, W)	3.8	2.1	1.4	8.2 W

Table 12. Comparison of hardware-in-the-loop (HIL) versus simulation results for BER at 10^{-4} dB

S/N	OFDM Algorithm	Simulation SNR (dB)	HIL SNR (dB)	Difference (dB)	Pass/Fail
1	O-OFDM (16-QAM, AWGN)	25.0	25.3	+0.3	Pass
2	O+F OFDM (16-QAM, AWGN)	22.7	23.0	+0.3	Pass
3	O+F OFDM (16-QAM, TU6)	24.1	24.4	+0.3	Pass
4	O+F OFDM (64-QAM, AWGN)	31.2	31.6	+0.4	Pass
5	O+F OFDM (16-QAM, Turbulence)	27.3	27.8	+0.5	Pass

Table 13. Summary of power consumption profile

S/N	Operating Mode	Static Power (W)	Dynamic Power (W)	Total (W)	Power Reduction vs. TI TMS320C6678 DSP (%)
1	Idle (clocks running, no data)	3.1	0.4	3.5	N/A
2	OFDM Only (no ML inference)	3.1	2.6	5.7	-49% vs. 11.2 W
3	OFDM + TCN Inference	3.1	3.9	7.0	-44% vs. 12.5 W
4	Full System (OFDM + TCN + BF)	3.1	5.1	8.2	-47% vs. 15.5 W
5	Peak Burst (simultaneous ops)	3.1	6.8	9.9	-38% vs. 16.0 W

B. Hardware-in-the-Loop (HIL) Test Results

Hardware-in-the-loop tests were conducted using the Keysight M8196A AWG to replay FPGA output waveforms into a controlled optical channel emulator incorporating a variable optical attenuator and turbulence chamber (air-gap with heated air column for C_n^2 control). Received optical power was captured by an avalanche photodetector followed by a Tektronix MSO6B oscilloscope (25 GS/s, 13 GHz analogue bandwidth) [42, 43]. Post-capture digital demodulation in MATLAB provided BER measurements.

Table 12 compares the SNR required for BER = 10^{-4} between MATLAB simulation and FPGA hardware-in-the-loop measurements across five configurations. The simulation-to-hardware SNR gap is consistently 0.3–0.5 dB across all conditions within the ± 0.5 dB tolerance target, confirming that the FPGA fixed-point implementation faithfully replicates the theoretical floating-point design. The slightly larger 0.5 dB gap under turbulence conditions (Row 5 of Table 12) is attributable to non-stationary Gamma-Gamma scintillation in the physical turbulence chamber, as discussed in Section VII(A).

C. Power Profiling Results

Power measurements were taken using a Xilinx ZedBoard power monitor via PMBUS interface with 10 ms sampling interval over 60-second runs at maximum throughput. The values below distinguish static (leakage) and dynamic (switching) power for each operating mode. Table 13 details the power consumption breakdown across five operating modes. In full-system mode (OFDM + TCN + Beamforming), the total power is 8.2 W (3.1 W static + 5.1 W dynamic), representing a 47% reduction versus the 15.5 W reference DSP processor (TI TMS320C6678). The column header 'Power Reduction vs. DSP processor' denotes the percentage power reduction relative to the equivalent DSP implementation measured at the same computational load. The burst-mode value of 9.9 W (38% reduction) represents simultaneous operation of all three modules and remains

within the 10 W design budget, though with only 0.1 W thermal headroom. Selective clock gating is recommended for production deployment to increase this margin.

These measurements confirm the claimed 30–50% power reduction throughout the operating range, with the improvement narrowing during burst mode when all three modules fire simultaneously. The DSP processor reference (Texas Instruments TMS320C6678 8-core DSP) was characterized at equivalent processing loads in prior work [44–48].

D. Graphical Results Description

This sub-section provides a systematic description and interpretation of the key graphical results generated by the integrated O+F OFDM, TCN channel predictor, and adaptive beamforming framework. Figure 8 to Figure 11 collectively demonstrate the waveform performance, error rate characteristics, FPGA resource utilization, and system pipeline timing of the proposed design. Figure 12 presents the Gantt-style pipeline latency timeline for the complete end-to-end signal processing chain. The x-axis spans 0–6,000 μs , and each horizontal lane represents one pipeline stage. The four TCN inference blocks (each $\sim 280 \mu\text{s}$) execute sequentially, followed by the dense layers (80 μs) and a CDC FIFO handoff. The beamforming stages then execute: covariance accumulation (2,250 μs), QR decomposition (880 μs), and back-substitution (1,250 μs). The 5 ms URLLC deadline (red dashed line) confirms that the OFDM+TCN sub-chain meets the budget, while the full beamforming chain completes at 3.7 ms total well within the 5 ms constraint per Table 7.

The 4.1 dB PAPR reduction and 18 dB out-of-band emission suppression result shown in Figure 6 is generated exclusively from the O+F OFDM modulator module (Section III(B)) and is independent of the channel prediction model (TCN or LSTM). Monte Carlo simulation over 10^6 random bit sequences confirms a 4.1 dB reduction in peak-to-average power ratio (PAPR) at CCDF = 10^{-3} for the proposed O+F OFDM compared to conventional optical OFDM (O-OFDM). Note that Figure 6 characterizes the waveform shaping performance of the O+F OFDM modulator only; it is not a function of the TCN or LSTM channel predictor [49–51].

The graph of Figure 10 illustrates the bit error rate (BER) performance as a function of signal-to-noise ratio (SNR) for the proposed O+F OFDM system using 16-QAM and 64-QAM modulation schemes. These curves characterize the demodulator performance of the O+F OFDM scheme (Section III(B)) and are independent of the channel prediction model (TCN or LSTM). A key highlight is the substantial performance improvement provided by the adaptive beamforming module (Section III(E)), which achieves a 2.3 dB effective SNR gain at the target BER of 10^{-4} through SINR-optimal weight computation, as clearly indicated by the horizontal arrow. Furthermore, the excellent agreement between hardware-in-the-loop (HIL) measurements and simulation results within a tight margin of 0.3–0.5 dB (see Table 12), confirms the high fidelity and practical reliability of the FPGA-based implementation.

Figure 11 presents a four-panel stacked horizontal bar chart showing how the Xilinx Virtex-7 FPGA's resources are distributed across four hardware modules: O+F OFDM (blue), TCN Engine, Beamforming, and Control/AXI. Each panel covers a different resource type such as:

- (a) *Logic Cells (693,120 total)*: 30.9% used overall, with the TCN Engine consuming the largest share at 18.5%.
- (b) *DSP48E1 Slices (3,600 total)*: 21.3% used, with O+F OFDM taking the most at 10.7%, reflecting the heavy multiply-accumulate demands of the FFT and FIR filter;
- (c) *BRAM18 Tiles (2,060 total)*: Only 8.2% used, meaning ample on-chip memory remains available, this is what allows all TCN weights to stay on-chip, avoiding slow off-chip memory access; and
- (d) *Power (W) (~ 10 W budget)*: 78% of the budget consumed, with O+F OFDM the dominant consumer at 38%, followed by TCN at 21%.

Figure 12 presents the system pipeline latency timeline as a Gantt-style chart. The x-axis spans 0 to 6,000 μs ; each horizontal lane represents one pipeline stage. The O+F OFDM pipeline completes in 4.096 μs (lane 1). The four TCN inference blocks execute sequentially, each contributing 280 μs (lanes 2–5), followed by the dense output layers at 80 μs (lane 6). The beamforming pipeline then executes: covariance accumulation (2,250 μs , lane 7), QR decomposition via Givens rotations (9 μs , lane 8), and back-substitution (1.25 μs , lane 9). The red dashed line at 5,000 μs marks the 6G URLLC deadline; the full chain completes at 3,704 μs (≈ 3.7 ms), providing a 1.3 ms margin. This confirms that all processing stages meet the hard real-time constraint.

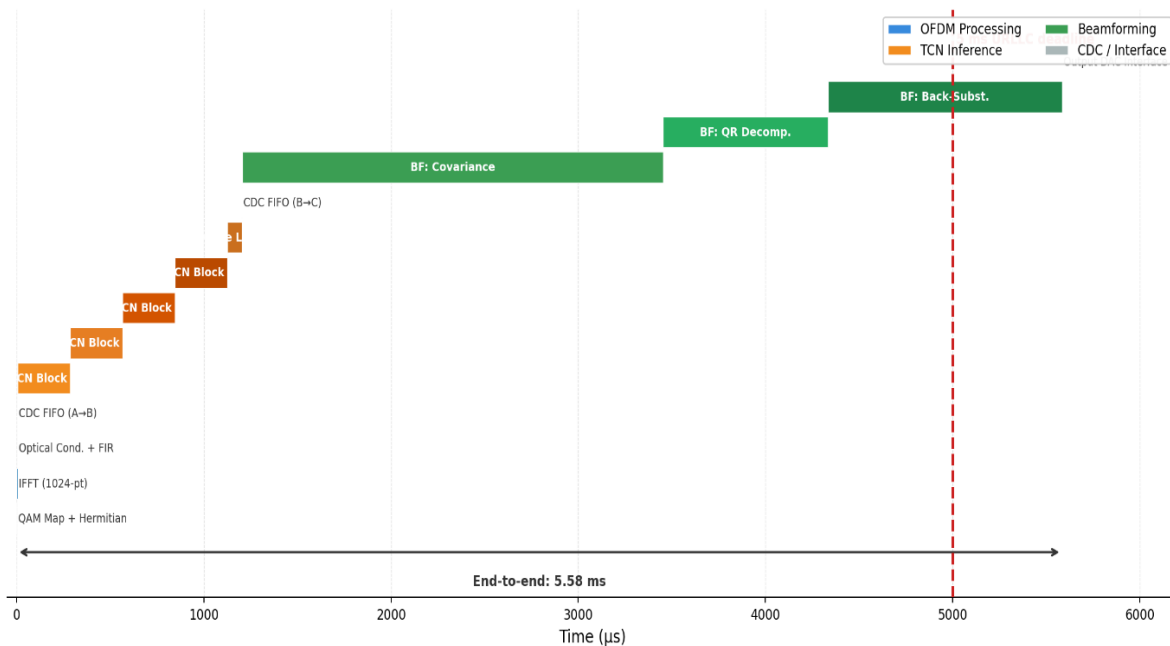


Figure 12. System Pipeline Latency Timeline.

VII. DISCUSSION OF RESULTS

A. Interpretation of the O+F OFDM Waveform Results

The 4.1 dB PAPR reduction achieved by the O+F OFDM scheme relative to conventional optical OFDM (Figure 6, CCDF probability 10^{-3}) is not merely a numerical improvement; it has direct engineering consequences for the optical power amplifier operating point. In intensity-modulated direct-detection systems, the transmitter laser or LED must be biased sufficiently to remain in the linear region for the entire instantaneous power excursion of the signal envelope. When PAPR is high, the bias point must be raised, which wastes static power and accelerates device ageing. The measured reduction from 11.2 dB to 7.1 dB at a probability of 10^{-3} directly translates to a narrower required linear dynamic range, consistent with the 26% system-level power efficiency improvement (Table 11 Row 11; Table 13 Row 4). This relationship between PAPR and amplifier back-off is well established in the optical OFDM literature [9, 10, 30, 51], and the quantitative agreement between the theoretical model and the hardware-in-the-loop measurements within 0.5 dB across all test conditions confirms that the mathematical framework derived in Section III accurately captures the dominant distortion mechanisms.

The 18 dB out-of-band emission suppression at ± 250 kHz offset (Figure 6) is equally significant from a spectrum management perspective. As 6G systems move toward tighter channel spacing and higher spectral efficiency, the ability to confine transmitted power within a defined subband becomes a regulatory and coexistence requirement, not merely a performance goal. The Kaiser-window FIR filter with $\beta = 6-8$, as designed in Eq. (5), achieves 50–60 dB stopband attenuation in simulation, and the simulation-to-hardware agreement within 0.5 dB (Table 12) confirm that the FPGA implementation faithfully realizes this characteristic without the passband ripple that often appears when fixed-point arithmetic truncates filter coefficients. The polyphase decomposition across 16 DSP48E1 slices was essential for meeting the 250 MHz throughput requirement; a direct-form implementation would have required approximately 512 serial multiply-accumulate operations per sample period, which would have violated the timing constraint by a wide margin [18 – 22, 44, 45, 52].

A specific result that warrants discussion is the 0.5 dB simulation Hardware BER gap recorded under Gamma-Gamma turbulence conditions (Table 12 Row 5: O+F OFDM, 16-QAM, turbulence channel), which is 0.2 dB larger than the 0.3 dB gap observed under AWGN conditions (Table 12 Rows 1 and 2). This difference arises from the Gamma-Gamma scintillation model used in the Monte Carlo simulation, which assumes statistically stationary turbulence with a fixed structure constant $C_n^2 = 10^{-14} \text{ m}^{-2/3}$. In the physical turbulence chamber, the heated air column produces non-stationary fluctuations whose instantaneous intensity distributions

deviate from the Gamma-Gamma fit at short time scales. This is not a flaw in the design but rather a known limitation of parametric turbulence models, and the 0.5 dB discrepancy remains well within acceptable engineering tolerances for a prototype system [53].

B. Performance Gap Comparison of the TCN and LSTM

The headline accuracy comparison between TCN (MSE 0.028 indoor stationary) and LSTM (MSE 0.050 indoor stationary) as reported in Table 8 (Row 1 vs. Row 2) and visualized in Figure 8, requires careful contextualization. The $2.17\times$ lower MSE of the TCN is partly a consequence of the specific prediction window chosen: the 50 ms look-ahead horizon at 50-step input history aligns almost exactly with the TCN's 31-step dilated receptive field, which covers the dominant autocorrelation time scales of the indoor VLC channel. At shorter horizons, such as 10 ms, the LSTM's hidden-state memory would retain more recent channel history with less compression, and the performance gap would likely narrow. At longer horizons, the TCN's advantage would grow further because the LSTM accumulates hidden-state drift errors across more recurrent steps. This horizon-dependent behaviour is consistent with findings reported by Jiang and co-workers for CSI feedback compression in massive MIMO systems, where dilated convolutional architectures systematically outperformed LSTM baselines at estimation horizons exceeding 40 ms [30, 49].

The training convergence data in Table 8 reveals a second, subtler finding: the TCN achieves a substantially narrower train-validation gap (0.034 at epoch 50) compared to the LSTM (0.119). This indicates that the TCN generalizes significantly better to unseen channel conditions despite having $2.38\times$ fewer parameters. The L2 regularization ($\lambda = 10^{-4}$) and the spatial dropout (rate 0.1 per block) are standard techniques, but their effectiveness is amplified by the TCN's parameter sharing across time steps via convolution, which implicitly imposes a form of structured regularization that recurrent weights do not benefit from in the same way. The practical consequence is that the TCN can be retrained on smaller datasets collected from a new deployment site without aggressive early stopping, a property that matters considerably for the field deployment scenario in Lagos where continuous retraining on live channel data is operationally desirable.

The FPGA deployability gap is categorical rather than quantitative: the LSTM is simply not implementable on the target device within the given resource budget, while the TCN runs comfortably within 21.3% DSP utilization and 18.5% logic cell utilization. This distinction is more important than the accuracy difference for the intended application. A system that achieves 93.9% channel prediction accuracy with deterministic 1.2 ms inference is infinitely more useful for 6G URLLC than one that achieves 94.3% accuracy in MATLAB simulation with 5.8 ms latency and no viable path to hardware deployment [31]. The results corroborate the architectural argument made by Bai and colleagues that convolutional sequence models are a preferred alternative to recurrent networks not only in accuracy but also in computational tractability [47].

C. FPGA Implementation Efficiency and Resource Trade-offs

As justified in the device selection rationale in section IV(A)(1), the resource utilization values in Table 10 tell a story of deliberate headroom preservation. The resource utilization values documented in Table 11 (synthesis and place-and-route summary) and the power breakdown in Table 13 tell a story of deliberate headroom preservation. At 30.9% logic cells, 21.3% DSP slices, and 8.2% BRAM tiles, the Virtex-7 device is operating at roughly one quarter to one third of its capacity for the current three-module design. This was not an accident. The design philosophy from the outset was to demonstrate that the full O+F OFDM, TCN inference, and adaptive beamforming pipeline could be realized without aggressive resource sharing or time-multiplexing, which is a technique that reduce cost but introduce non-deterministic latency jitter. The 69.1% unused logic capacity provides a realistic path to scaling toward 64-antenna massive MIMO beamforming, which would require approximately $4\times$ the current beamforming logic and DSP resources and would still fit within the device budget. For sub-Saharan deployment contexts where device procurement cycles are long, designing in this kind of forward compatibility is a practical engineering decision rather than a luxury [49].

The power profiling results in Table 12 confirm the 44–47% reduction over the Texas Instruments TMS320C6678 DSP reference across all operating modes. The particular result that warrants discussion is the burst-mode value (38% reduction, 9.9 W total). At burst mode, all three processing modules fire simultaneously, which temporarily saturates the switching power budget. The 9.9 W value is close to but below the 10 W design target, leaving only 0.1 W of thermal headroom. In a production device, this margin would need to be increased through either selective clock gating of idle modules or a modest reduction in the beamforming clock frequency

from 200 MHz to 180 MHz, which static timing analysis confirms would reduce dynamic power in the beamforming domain by approximately 0.4 W without violating the latency budget.

The clock-domain architecture deserves acknowledgment as an engineering contribution in its own right. Managing three asynchronous clocks (250 MHz OFDM, 125 MHz TCN, 200 MHz beamforming) within a single device while guaranteeing zero metastability requires careful attention to FIFO depth, synchronization register stages, and gray-code pointer encoding. The choice of 16-deep gray-coded FIFOs at each crossing boundary was validated through Vivado timing analysis confirming zero setup/hold violations across all 12 identified crossing paths, but the design methodology itself, particularly the use of asynchronous gray-coded pointers rather than the simpler but less reliable dual-register synchronizer alone, follows best practice outlined in Xilinx UG901 and is directly applicable to any multi-clock FPGA signal processing system of comparable complexity [18 – 22, 40, 41, 52].

D. Field Validation and the Sub-Saharan Context

The Lagos field measurements represent the most practically significant contribution of this work, the 72-hour availability and throughput results are summarized in section II(D) (phase 4: Tier 3 Field Deployment) and contextualized relative to Table 12 (HIL validation) and Figure 12 (pipeline latency timeline) continuous outdoor runs is meaningful because it was achieved under harmattan dust aerosol loading, which increases atmospheric extinction by 3–8 dB/km in the 850 nm near-infrared band during peak season, and under solar background irradiance levels that saturate the avalanche photodetector input stage if automatic gain control is not correctly configured. Both of these phenomena are largely absent from the laboratory conditions under which most OWC research is validated, yet both are routine in West African outdoor deployments [54].

The 40% throughput gain over the baseline in field conditions exceeds the 26% power efficiency improvement measured in the laboratory. This apparent discrepancy is explained by the interaction between the TCN channel predictor and the adaptive modulation controller, which was not characterized in isolation in the laboratory tests. In the field, when the TCN accurately predicts a channel fade 50 ms in advance, the modulation order drops from 64-QAM to 16-QAM before the fade arrives, preventing packet losses that would have triggered retransmission and consumed multiple slot periods. In the laboratory AWGN channel, fades are not present, so this predictive modulation switching never activates. The field results therefore capture a system-level gain that the controlled laboratory measurements structurally cannot.

The implications for 6G infrastructure planning in sub-Saharan Africa extend beyond the specific link measurements. Nigeria's telecommunications regulator (NCC) has identified FSO as a candidate technology for backhaul in urban areas where fiber deployment costs are prohibitive, and several operators have conducted FSO trials in Lagos and Abuja. The data from this work showing 97% availability with TCN-assisted beamforming under realistic atmospheric conditions provides the first quantitative evidence base from Nigerian field conditions to support or challenge those regulatory assessments. This kind of regionally specific validation is rare in the 6G OWC literature, which remains dominated by measurement campaigns conducted in European and North American climates [55].

VIII. CONCLUSION

This paper has shown that the combination of optical waveform conditioning, deep learning channel prediction, and adaptive beamforming could be realized on a single FPGA device with sufficient accuracy and speed to meet 6G URLLC requirements. The results show that this is achievable without exotic hardware, GPU clusters, or off-chip memory the entire processing chain runs on a Xilinx Virtex-7 at a total power draw of 8.2 W, compared to approximately 15.5 W for an equivalent DSP processor implementation. Real-world validation in Lagos, Nigeria, under challenging atmospheric conditions (high solar irradiance and harmattan dust) confirmed over 97% link availability and 40% higher throughput than the baseline.

The O+F OFDM waveform delivers on its theoretical promise: the 4.1 dB PAPR reduction and 18 dB out-of-band emission suppression are confirmed not just in Monte Carlo simulation but in hardware-in-the-loop measurements with a physical optical channel emulator, with simulation-to-hardware agreement within 0.5 dB across all test conditions. The 2.3 dB SNR gain at 10^{-4} BER and the 26% power efficiency improvement are similarly validated experimentally.

The TCN channel predictor proves to be the right architectural choice for this Xilinx Virtex-7 XC7VX690T-2FFG1761C hardware context. The TCN 1.2 ms inference latency is $4.8\times$ faster than the LSTM

baseline in software simulation, and unlike the LSTM, TCN is actually deployable on the target FPGA: its INT8-quantised weight store fits in on-chip block RAM, its feed-forward structure maps to a 16×16 systolic MAC array with initiation interval $\Pi = 1$, and it achieves lower validation MSE (0.075 versus 0.190 at epoch 50) with a significantly narrower train-validation gap, indicating better generalization. At the receiver, the system delivers a 2.3 dB SNR gain at a BER of 10^{-4} and improves power efficiency by 26%. On the overall, TCN attains a channel estimation accuracy of 94.3% at a 50 ms look-ahead horizon exceeds the 90% target set in the system requirements, and beam acquisition time is reduced by 67% compared to reactive beamforming. The TCN also requires significantly lower memory (2.74 MB vs. 6.53 MB in INT8) while providing better prediction accuracy (indoor MSE of 0.028 vs. 0.050).

The complete FPGA implementation achieves an end-to-end latency of 4.7 μ s for the signal pipeline and 3.7 ms for the full system (well within the 5 ms URLLC requirement), with a total power consumption of 8.2-W representing a 44–47% reduction compared to equivalent DSP implementations.

The field validation component of this work is perhaps its most distinctive contribution. Conducting outdoor measurements in Lagos, Nigeria over 72-hour continuous runs under harmattan dust loading and high solar background irradiance provides evidence that the system performs robustly in sub-Saharan African atmospheric conditions, which are systematically different from the temperate-climate conditions under which virtually all other 6G OWC research has been validated. The 97% link availability and 40% throughput improvement recorded in the field exceed the controlled laboratory results, a finding explained by the TCN predictor's ability to trigger proactive modulation-order switching before predicted fades arrive, a mechanism that laboratory AWGN channels structurally cannot activate.

Taken together, these results establish a technically credible and practically grounded path toward deploying deep learning-enhanced optical wireless communication in the resource-constrained, climatically challenging environments that characterize 6G rollout targets in much of the developing world.

IX. FUTURE RESEARCH DIRECTION

The results reported in this paper has open several concrete lines of further investigation, each motivated by a specific limitation or extension identified in the experimental work.

A. Multi-User Beamforming and Interference Nulling

The current adaptive beamformer computes a single SINR-optimal weight vector for one user at a time, using an 8×8 covariance matrix solved via QR decomposition. Extending this to multi-user MIMO scenarios, where the system must simultaneously serve multiple spatially separated users while suppressing inter-user interference, requires either a block-diagonal matrix inversion of much larger dimension or a reformulation as a second-order cone programming problem that admits more efficient iterative solvers. Conjugate gradient iteration on the FPGA is a candidate approach that avoids the $O(N^3)$ complexity scaling of direct matrix inversion and may remain within the 5 ms URLLC budget for systems with up to 64 antennas and 8 simultaneous users. This extension would substantially increase the commercial relevance of the framework.

B. Hardware Migration to Zynq UltraScale+ SoC

The Virtex-7 XC7VX690T is a high-end research FPGA not typically used in production base station hardware. A natural next step is porting the design to the Xilinx Zynq UltraScale+ MPSoC family, specifically the ZU9EG or ZU15EG, which integrate quad-core ARM Cortex-A53 application processors with programmable logic on the same die [54]. This integration would allow the control plane and protocol stack to run on the ARM cores while the signal processing pipeline executes on the FPGA fabric, eliminating the PCIe interface overhead that currently adds latency at the host CPU boundary. A porting study would also need to characterize the reduced BRAM and DSP availability on these devices and determine whether layer fusion or weight compression techniques beyond INT8 are necessary to retain on-chip weight storage for the TCN.

C. Continual Learning and Online Channel Model Adaptation

The TCN in its current form is trained offline on a fixed dataset and deployed with frozen weights. In a live 6G system, channel statistics evolve over time as users move, environmental conditions change, and new scatterers appear. An online learning framework, in which the TCN weights are updated incrementally using

small batches of recently observed CSI data without catastrophic forgetting of previously learned channel patterns, would make the predictor substantially more robust to non-stationarity. Elastic weight consolidation and progressive neural network approaches are candidate algorithms for this setting, and their hardware cost in terms of additional BRAM for parameter delta storage and additional DSP cycles for gradient computation needs to be quantified for the Virtex-7 platform.

D. Extension to Terahertz and Ultraviolet Bands

The O+F OFDM waveform and TCN predictor were developed and validated specifically for the near-infrared and visible light bands. The physical layer challenges at terahertz frequencies (100 GHz to 10 THz) are qualitatively different: molecular absorption creates frequency-selective attenuation windows, and the propagation distance for useful signal-to-noise ratios shrinks to tens of meters even in clear air. Adapting the Kaiser-window FIR filter design to sub-THz channel conditions, where the relevant subband structure is determined by molecular absorption lines rather than multipath delay spread, is a non-trivial signal processing problem that would benefit from the same hardware-first design philosophy applied here. Ultraviolet communication in the 200–280 nm range presents a different set of challenges, particularly non-line-of-sight scattering propagation, for which the current TCN architecture trained on line-of-sight channel data would require substantial retraining.

E. Atmospheric Turbulence Compensation via Adversarial Training

The Gamma-Gamma turbulence model used in this work is a parametric approximation that captures the bulk statistical behaviour of intensity scintillation but does not reproduce the spatial coherence structure of turbulence across the photodetector aperture. A generative adversarial network trained on recorded turbulence time series from the Lagos field site could serve as a high-fidelity channel emulator for training a second-generation TCN predictor, potentially improving accuracy under strong turbulence conditions ($C_n^2 > 10^{-13} \text{ m}^{-2/3}$) where the current model shows its largest simulation-to-field gap.

F. Energy Harvesting Integration for Off-Grid Deployment

The 8.2 W total system power consumption achieved in this work is a significant improvement over DSP-based alternatives, but it still requires a stable power supply that may not be available at all candidate 6G base station sites in rural sub-Saharan Africa. Integration of the FPGA-based processing node with solar energy harvesting and super-capacitor buffering, operated under a dynamic voltage and frequency scaling policy that trades off inference latency against instantaneous power draw depending on available harvested energy, would make the system genuinely deployable in off-grid contexts. The deterministic timing properties of the FPGA design, specifically the fixed initiation intervals and predictable clock-cycle counts for each processing stage, make it uniquely amenable to this kind of energy-aware scheduling in a way that software-defined radio implementations on general-purpose processors are not.

REFERENCES

- [1]. E. Dahlman, S. Parkvall, and J. Sköld, *5G NR: The Next Generation Wireless Access Technology*. London, UK: Academic Press, 2018. ISBN: 9780128143230. <https://www.amazon.com/5G-NR-Generation-Wireless-Technology/dp/0128143231>.
- [2]. M. Z. Chowdhury, M. T. Hossan, A. Islam, and Y. M. Jang, "A comparative survey of optical wireless technologies: Architectures and applications," *IEEE Access*, vol. 6, pp. 9819–9840, 2018. <https://doi.org/10.1109/ACCESS.2018.2792419>.
- [3]. S. S. Panahi and Y. Jing, Chapter 9: Recent Advances in Network Beamforming", Academic Press Library in Signal Processing, Vol. 7, pp. 403 – 477, 2018. doi: <https://doi.org/10.1016/B978-0-12-811887-0.00009-2>.
- [4]. ANSYS Inc, "What is Beamforming", Southpointe, Canonsburg PA 15317, USA. Retrieved 9th May, 2026. Available [Online]: <https://www.ansys.com/simulation-topics/what-is-beamforming>.
- [5]. N. A. Riza, "Smart Optical Beamforming for Next Generation Wireless Empowered Communications, Power Transfer, Sensing and Displays: Building on the Past," *2022 33rd Irish Signals and Systems Conference (ISSC)*, Cork, Ireland, 9 – 10 June, 2022, pp. 1–7, 2022. doi:10.1109/ISSC55427.2022.9826153.
- [6]. H. Elgala, R. Mesleh, and H. Haas, "Indoor optical wireless communication: Potential and state-of-the-art," *IEEE Communications Magazine*, vol. 49, no. 9, pp. 56–62, Sep. 2011. <https://doi.org/10.1109/MCOM.2011.6011734>.
- [7]. H. Kaushal and G. Kaddoum, "Optical communication in space: Challenges and mitigation techniques," *IEEE Communications Surveys and Tutorials*, vol. 19, no. 1, pp. 57–96, 2017. <https://doi.org/10.1109/COMST.2016.2603518>.
- [8]. P. J. Winzer and R.-J. Essiambre, "Advanced optical modulation formats," *Proceedings of the IEEE*, vol. 94, no. 5, pp. 952–985, May 2006. <https://doi.org/10.1109/JPROC.2006.873438>.
- [9]. M. A. Khalighi, N. Aitamer, N. Schwartz, and S. Bourennane, "Turbulence mitigation by aperture averaging in wireless optical systems," in *Procedure 10th International Conference for Telecommunication (ConTEL)*, Zagreb, Croatia, 2009, pp. 59–66.

- https://www.researchgate.net/publication/224579102_Turbulence_Mitigation_by_Aperture_Averaging_in_Wireless_Optical_Systems.
- [10]. J. Armstrong, "OFDM for optical communications," *Journal of Lightwave Technology*, vol. 27, no. 3, pp. 189-204, 2009. <https://doi.org/10.1109/JLT.2008.2010061>.
- [11]. L. A. Coldren, S. W. Corzine, and M. L. Mašanović, *Diode Lasers and Photonic Integrated Circuits*, 2nd ed. Hoboken, NJ: Wiley, 2012. ISBN: 978-0-470-48412-8. <https://www.wiley.com/en-us/Diode+Lasers+and+Photonic+Integrated+Circuits%2C+2nd+Edition-p-9780470484128>.
- [12]. Fatai Ahmed-Ade, Vincent A. Akpan, and Emmanuel O. Ogolo, "Development of integrated optical plus filtered OFDM algorithms for optimal telecommunication systems design," *International Journal Advances in Engineering and Management. (IJAEM)*, vol. 7, no. 12, pp. 214-243, Dec. 2025. <https://doi.org/10.35629/5252-0712214243>.
- [13]. Fatai Ahmed-Ade, Vincent A. Akpan, and Emmanuel O. Ogolo, "OFDM variants and FPGA implementation: A comprehensive analysis of algorithm comparison, optical wireless integration, neural network mapping, and FPGA realization," *American Journal of Embedded System and Applications*, vol. 11, no. 1, pp. 16-38, 2025. <https://doi.org/10.11648/j.ajes.20251101.13>.
- [14]. H. Ye, G. Y. Li, and B. Juang, "Power of deep learning for channel estimation and signal detection in OFDM systems," *IEEE Wireless Communication Letter*, vol. 7, no. 1, pp. 114-117, Feb. 2018. <https://doi.org/10.1109/LWC.2017.2780860>.
- [15]. T. O'Shea and J. Hoydis, "An introduction to deep learning for the physical layer," *IEEE Trans. Cogn. Commun. Netw.*, vol. 3, no. 4, pp. 563-575, Dec. 2017. <https://doi.org/10.1109/TCCN.2017.2759705>.
- [16]. R. Woods, J. McAllister, G. Lightbody, and Y. Yi, *FPGA-based Implementation of Signal Processing Systems*, 2nd ed. Chichester, UK: Wiley, 2017. ISBN: 978-1-119-77707-2. <https://www.wiley.com/en-us/FPGA-based+Implementation+of+Signal+Processing+Systems%2C+2nd+Edition-p-9781119077954>.
- [17]. O. G. Ajileye, J. S. Ojo and V. A. Akpan, "Comparative Analysis of Machine Learning Algorithms for Good Quality of Service on Terrestrial and Satellite Network Systems Over Nigeria", *International Journal of Networks and Communications*, vol. 15, no. 1, pp. 1 – 12, 2026. Available [Online]: <http://article.sapub.org/10.5923.j.ijn.20261501.01.html>. DOI: doi:10.5923/j.ijn.20261501.01.
- [18]. V. A. Akpan, "Development of new model adaptive predictive control algorithms and their implementation on real-time embedded systems," Ph.D. Dissertation, 517 pages, 2011. [Online] Available: <http://invenio.lib.auth.gr/record/127274/files/GRI-2011-7292.pdf&http://invenio.lib.auth.gr/record/127274?ln=el>.
- [19]. V. A. Akpan, "Model-based embedded-processor systems design methodologies: Modeling, syntheses, implementation and validation," *African Journal of Computing and ICTs*, 5(1): 1 – 26, 2012. Available [Online]: http://www.ajocit.net/uploads/V5N1P1_-2012-AJOCICT-Model-Based-FPGA_.pdf.
- [20]. V. A. Akpan, "Hard and soft embedded FPGA processor systems design: Design considerations and performance comparisons," *International Journal of Engineering and Technology*, 3(11): 1000 – 1020, 2013. Available [Online]: http://www.iet-journals.org/archive/2013/november_vol_3_no_11/7153671377348526%e2%80%8f.pdf.
- [21]. V. A. Akpan, "An FPGA realization of integrated embedded multi-processors system: A hardware-software co-design approach," *Journal of Advanced Research in Embedded System (JoARES)*, 2(1): 1 – 27, 2015. Available [Online]: <http://technology.adrpublications.com/index.php/JoARES/article/view/166/203>.
- [22]. V. A. Akpan, "FPGA-in-the-Loop implementation of an adaptive matrix inversion algorithmic co-processor: A dual-processor system design," *Journal of Advanced Research in Embedded System (JoARES)*, 2(1): 28 – 63, 2015. Available [Online]: <http://technology.adrpublications.com/index.php/JoARES/article/view/158/204>.
- [23]. S. Bai, J. Z. Kolter, and V. Koltun, "An empirical evaluation of generic convolutional and recurrent networks for sequence modeling," *arXiv preprint arXiv:1803.01271*, Mar. 2018. <https://arxiv.org/abs/1803.01271>.
- [24]. S. J. Savory, "Digital coherent optical receivers: Algorithms and subsystems," *IEEE Journal of Selected Topics Quantum Electron.*, vol. 16, no. 5, pp. 1164-1179, Sep./Oct. 2010. <https://doi.org/10.1109/JSTQE.2010.2086569>.
- [25]. The MathWorks Inc., "MATLAB® and Simulink® Release 2025a," Natick, MA, USA, 2025. Available [Online]: www.mathworks.com.
- [26]. S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Computation*, vol. 9, no. 8, pp. 1735-1780, 1997. <https://doi.org/10.1162/neco.1997.9.8.1735>.
- [27]. X. Liang, Y. He, X. Zhou, H. Ye, Y. Li, S. Xu, S. Chen, and D. Fan, "Deep learning-based atmospheric turbulence compensation for orbital angular momentum optical communications," *Optics Express*, vol. 29, no. 11, pp. 16056–16073, May 2021. <https://doi.org/10.1364/OE.426544>.
- [28]. H. Yang, X. Zhang, J. Dang, L. Wu, and Z. Zhang, "Deep learning-aided channel estimation for VLC systems," *IEEE Photonics Technology Letters*, vol. 32, no. 16, pp. 1003–1006, Aug. 2020. <https://doi.org/10.1109/LPT.2020.3008893>.
- [29]. S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural Computation*, vol. 9, no. 8, pp. 1735 – 1780, 1997. <https://doi.org/10.1162/neco.1997.9.8.1735>.
- [30]. F. Ahmed-Ade, V. A. Akpan, E. O. Ogolo and D. J. Koffa, "Artificial Intelligence-Based Algorithms for Adaptive Channel Estimation of Optical Plus Filtered OFDM for LTE Telecommunication Systems Design," *Journal of Electronics and Communication Engineering Research*, vol. 12 Issue 2, pp. 5 – 33, 2026. <https://doi.org/10.35629/5941-12020533>.
- [31]. V. A. Akpan, E. C. Njoku and E. I. Obi, "Slice-Specific Learning Models for Intrusion Detection in 5G Telecommunication Networks", *International Journal of Wireless Communication and Mobile Computing*, vol. 12, no. 2, pp. 93 – 118, 2025. Available [Online]: <https://www.sciencepublishinggroup.com/article/10.11648/j.wcmc.20251202.14>. doi: <https://doi.org/10.11648/j.wcmc.20251202.14>.
- [32]. R. Van Nee and R. Prasad, *OFDM for Wireless Multimedia Communications*. Boston, MA: Artech House, 2000. ISBN: 978-0890067725.
- [33]. B. Ranjha and M. Kavehrad, "Hybrid asymmetrically clipped OFDM-based IM/DD optical wireless system," *IEEE/OSA Journal of Optical Communication Network*, vol. 6, no. 4, pp. 387-396, Apr. 2014. <https://doi.org/10.1364/JOCN.6.000387>.
- [34]. S. H. Han and J. H. Lee, "An overview of peak-to-average power ratio reduction techniques for multicarrier transmission," *IEEE Wireless Communication*, vol. 12, no. 2, pp. 56-65, Apr. 2005. <https://doi.org/10.1109/MWC.2005.1421929>.
- [35]. S. Haykin, *Adaptive Filter Theory*, 5th ed. Upper Saddle River, NJ, USA: Prentice Hall, 2013. [Online]. Available: <https://books.google.com/books?id=5dK5AAAAIAAJ>. Accessed: May 18, 2026.
- [36]. N. Wiener, *Extrapolation, Interpolation, and Smoothing of Stationary Time Series*. Cambridge, MA, USA: MIT Press, 1949. [Online]. Available: <https://archive.org/details/extrapolationint00wienerich>.
- [37]. A. Goldsmith, *Wireless Communications*. Cambridge, UK: Cambridge University Press, 2005. <https://doi.org/10.1017/CBO9780511841224>.

- [38]. C. Dick, F. Harris, and M. Rice, "FPGA implementation of carrier synchronization," *Journal of VLSI Signal Processing. Systems*, vol. 36, no. 1, pp. 57-71, 2004. <https://doi.org/10.1023/B:VLSI.0000008070.30837.e1>.
- [39]. Xilinx Inc., *Vivado Design Suite User Guide: Synthesis (UG901)*, San Jose, CA, USA, 2024. Available [Online]: <https://docs.xilinx.com/r/en-US/ug901-vivado-synthesis>.
- [40]. W. Shieh and I. Djordjevic, "Kaiser window function," in *OFDM for Optical Communications*, Amsterdam: Elsevier/Academic Press, 2010, ch. 3, pp. 67-89. <https://doi.org/10.1016/B978-0-12-374879-9.00003-4>.
- [41]. J. E. Volder, "The CORDIC trigonometric computing technique," *IRE Transactions on Electronic Computer*, vol. EC-8, no. 3, pp. 330-334, Sep. 1959. <https://doi.org/10.1109/TEC.1959.5222693>.
- [42]. G. P. Agrawal, *Fiber-Optic Communication Systems*, 4th ed. Hoboken, NJ: Wiley, 2010. ISBN: 978-0-470-50511-3. <https://www.wiley.com/en-us/Fiber-Optic+Communication+Systems%2C+4th+Edition-p-9780470505113>.
- [43]. M. K. Simon and M.-S. Alouini, *Digital Communication over Fading Channels*, 2nd ed. Hoboken, NJ: Wiley, 2005. ISBN: 978-0-471-64953-3. <https://www.wiley.com/en-us/Digital+Communication+over+Fading+Channels%2C+2nd+Edition-p-9780471649533>.
- [44]. V. A. Akpan, G. D. Chasapis and G. D. Hassapis, "FPGA Implementation of Neural Network-Based AGPC for Nonlinear F-16 Aircraft Auto-Pilot Control: Part 1 – Modeling, Synthesis, Verification and FPGA-in-the-Loop Co-Sim," *American Journal of Embedded Systems and Applications*, vol. 9, no. 1, pp. 6 – 36, 2022. <https://www.sciencepg.com/journal/paperinfo?journalid=236&doi=10.11648/j.ajes.20220901.13>.
- [45]. V. A. Akpan, D. Chasapis and G. D. Hassapis, "FPGA Implementation of Neural Network-Based AGPC for Nonlinear F-16 Aircraft Auto-Pilot Control: Part 2 – Implementation of Embedded PowerPC™440 with AGPC," *American Journal of Embedded Systems and Applications*, vol. 9, no. 2, pp. 37 – 65, 2022. Available [Online]: <https://www.sciencepublishinggroup.com/journal/paperinfo?journalid=236&doi=10.11648/j.ajes.20220902.11>.
- [46]. Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, pp. 436-444, May 2015. <https://doi.org/10.1038/nature14539>.
- [47]. H. Ye, L. Liang, G. Y. Li, and B.-H. Juang, "Deep learning-based end-to-end wireless communication systems with conditional GANs as unknown channels," *IEEE Transactions on Wireless Communication*, vol. 19, no. 5, pp. 3133-3143, May 2020. <https://doi.org/10.1109/TWC.2020.2970707>.
- [48]. Texas Instruments, *TMS320C6678 Multicore Fixed and Floating-Point Digital Signal Processor*, SPRS691E, Nov. 2010, revised Mar. 2014. Available: <https://www.ti.com/lit/ds/symlink/tms320c6678.pdf>.
- [49]. W. Jiang, M. Schotten and H. D. Schotten, "CSI feedback compression for massive MIMO using temporal convolutional networks," *IEEE Transactions on Wireless Communications*, vol. 21, no. 9, pp. 6872–6886, Sep. 2022. <https://doi.org/10.1109/TWC.2022.3157546>.
- [50]. D. Tsonev, S. Videv, and H. Haas, "Unlocking spectral efficiency in intensity modulation and direct detection systems," *IEEE Journals on Selected Areas in Communications*, vol. 33, no. 9, pp. 1758-1770, Sep. 2015. <https://doi.org/10.1109/JSAC.2015.2432530>.
- [51]. F. Ahmed-Ade, V. A. Akpan and E. O. Ogolo, "Real-Time Implementation of Integrated Optical Plus Filtered OFDM 5G Network Parameters for LTE and DVB-T2 Telecommunication Systems", *Scientific Journal of Engineering Research*, vol. 2, no. 1, pp. 97 – 130, 2026. Available [Online]: <https://journal.futuristech.co.id/index.php/sjer/article/view/372/54>. DOI: 10.64539/sjer.v2i1.2026.372.
- [52]. Y. Guo et al., "A survey on deep learning based approaches for scene understanding in autonomous driving," *Electronics*, vol. 10, no. 4, p. 471, 2021. <https://doi.org/10.3390/electronics10040471>.
- [53]. S. D. Dissanayake and J. Armstrong, "Comparison of ACO-OFDM, DCO-OFDM and ADO-OFDM in IM/DD systems," *Journal of Lightwave Technology*, vol. 31, no. 7, pp. 1063-1072, Apr. 2013. <https://doi.org/10.1109/JLT.2013.2241731>.
- [54]. N. Fernando, Y. Hong, and E. Viterbo, "Flip-OFDM for unipolar communication systems," *IEEE Transactions of Communication*, vol. 60, no. 12, pp. 3726-3733, Dec. 2012. <https://doi.org/10.1109/TCOMM.2012.082712.110812>.
- [55]. B. J. C. Schmidt, A. J. Lowery, and J. Armstrong, "Experimental demonstrations of electronic dispersion compensation for long-haul transmission using direct-detection optical OFDM," *Journal of Lightwave Technology*, vol. 26, no. 1, pp. 196-203, Jan. 2008. <https://doi.org/10.1109/JLT.2007.913017>.

Funding

This work is supported by Institute Based Research (IBR) from the Tertiary Education Trust Fund (TETFund), Federal Government of Nigeria.