

A Hybrid Optical Wireless Network Architecture with an Efficient QoS-Provisioning Transmission Strategy for us-Scale Latency Data Centers

Georgios Drainakis, Peristera Baziana

Department of Informatics and Telecommunications, University of Thessaly, Greece
gdrainakis@uth.gr, baziana@uth.gr

Abstract—Driven by the emergence of high bandwidth, deterministic ultra-low latency cloud services, data center networks (DCNs) are increasingly challenged by highly heterogeneous traffic with diverse quality-of-service (QoS) requirements. Although optical networking has dominated DCN technology due to its capacity, most existing solutions remain largely application-agnostic. To address this gap, we propose a hybrid fiber-optical and free-space optical wireless DCN architecture along with medium access control (MAC) protocols that enable QoS provisioning through priority-based transmission. The architecture comprises two complementary domains: an intra-rack fiber-optical domain for server-to-server and server to Top-of-Rack (ToR) communication, under traffic priority constraints and an inter-rack free-space optical domain leveraging virtually imaged phased array (VIPA) technology to support prioritized low-latency transmission across multiple wavelengths within a flat, single-hop topology. Simulation results demonstrate bandwidth efficiency reaching 100% over a wide load range and 94% under nominal load. Moreover, our proposed solution enables strict application differentiation and sub-us end-to-end latency for high priority traffic, enabling support for emerging mobile-edge services, immersive media and real-time future use cases.

Index Terms—Free Space Optics, DC Networks, MAC Protocols

I. INTRODUCTION

Over the past decade, the rapid emergence of data-intensive and latency-sensitive applications has fundamentally reshaped the requirements imposed on communication networks. Artificial intelligence (AI) and Machine Learning (ML) workloads now dominate this landscape, driven by unprecedented growth in training and inference tasks across cloud, edge and industrial environments. Recent reports indicate that tens of billions of Internet-of-Things (IoT) devices collectively generate data volumes exceeding tens of zettabytes annually, with AI being indispensable for extracting value from this telemetry at scale [1]. Beyond AI, additional technology frontiers, e.g., quantum networking, fifth generation (5G) and beyond systems, immersive media applications such as augmented reality (AR), virtual reality (VR) and extended reality (XR) are converging to create heterogeneous traffic patterns with extreme bandwidth, latency and reliability demands [2].

To accommodate this surge in computational and networking demand, the global data center (DC) ecosystem is entering a phase of unprecedented expansion. Installed DC capacity is projected to nearly double by the end of the decade, with

AI workloads alone expected to consume approximately half of total capacity within a few years [3]. At the same time, the availability of suitable land, power delivery infrastructure and cooling capacity, particularly in metropolitan regions, is becoming increasingly limited and costly. These trends have elevated environmental sustainability to a primary design objective, prompting regulatory initiatives such as the European Union's Data Centre Energy Efficiency Package [3].

While traditional electrical interconnects historically offered cost-effective connectivity, their susceptibility to attenuation, crosstalk and electromagnetic interference makes them unsuitable for the high bandwidth densities and extended reach demanded by modern DCs [4]. As a result, optical networking has emerged as the dominant enabling technology for next-generation DC fabrics. Contemporary fiber-optic links routinely support data rates of 400 to 800 Gbps, with early deployments of 1.6 Tbps systems already underway [5]. Optical interconnects offer superior scalability, immunity to electromagnetic interference and improved thermal stability, making them essential for large-scale architectures. Despite their advantages, fiber-based optical DC networks face significant scalability challenges as DC size and traffic heterogeneity grow. Physically, the deployment and management of vast numbers of fiber cables introduce substantial complexity, leading to dense cable bundles, constrained airflow and increased operational overhead [6].

A promising technology to address the scalability, flexibility and energy limitations of fiber-based DC networks (DCNs) is Free-Space Optics (FSO), which enables optical wireless interconnects within DCNs. FSO links offer many of the performance advantages of fiber-optic systems, including high bandwidth, low latency and immunity to electromagnetic interference, but without the physical constraints of cabling [7], [8]. Optical wireless DCNs provide the ability to dynamically allocate capacity according to traffic demand, mitigating the limitations of fixed fiber topologies, lowering maintenance requirements and alleviating congestion within dense rack environments [9]. Recent experimental demonstrations have shown that FSO links can achieve bandwidths comparable to that of fiber-optic systems, while emerging technologies such as microcombs offer compact and power-efficient light sources suitable for next-generation optical wireless networks [7]. Despite these advances, a key research challenge remains:

integrating FSO components into existing DCN architectures and defining transmission mechanisms, e.g., medium access control (MAC) protocols, to efficiently support the wide diversity of applications and traffic patterns observed in modern DCN workloads [8].

The current state of the art (SotA) on optical wireless interconnects for DCNs remains limited. Most efforts have concentrated on the physical layer, focusing on high-speed optical wireless transmission, modulation formats and experimental demonstration of intra-rack or inter-rack links. For instance, in [9] a photonic integrated multicast switch (PIC-MCS) is proposed for dynamic optical packet switching across racks, demonstrating ns switching and error-free operation at 58 Gbps in a 4×4 rack setup. Similarly, in [10] a passive metasurface-based optical wireless interconnect is introduced, showing multicast capability with polarization control. Other contributions focus on link-level performance, including terabit spatial multiplexed Multiple Input Multiple Output (MIMO) optical wireless communications for inter-rack transmission [11], solid-state 360-degree beamforming for multicast links [12] and plug-and-play self-alignment systems achieving tens of Gbps at short distances [13]. While these studies advance the understanding of high-speed FSO and highlight the feasibility of optical wireless links, they largely neglect aspects such as traffic-aware scheduling and transmission coordination, network integration and key metrics such as end-to-end latency and DCN-level throughput. Only a few works address the challenges of scaling beyond a single rack or small cluster. To this end, a scalable matrix-ring FSO interconnect for large-scale DCNs is proposed in [14], demonstrating architectural concepts for thousands of nodes. The proposed ODRAD architecture stands as most relevant contribution to DCN-level FSO integration [15], which combines distributed software-defined networking (SDN), deep reinforcement learning (DRL), semiconductor optical amplifiers (SOAs) and arrayed waveguide grating routers (AWGRs). ODRAD simulates end-to-end (E2E) server-to-server latency at microsecond scales and shows scaling potential up to 40K servers, thereby directly addressing both the architectural and operational requirements of modern, heterogeneous DCNs.

To address the identified gaps, i.e., lack of traffic-aware coordination and network-level integration and scalability beyond small clusters, we propose a novel, scalable, hybrid optical DCN architecture that integrates high-performance fiber-optical interconnects within racks and optical wireless links between racks, forming a fully E2E optical DCN. On one hand, intra-rack communication leverages conventional fiber optics to achieve high throughput and low latency among servers within the same rack, as long as between servers and corresponding Top-of-Rack (ToR) switch. On the other hand, inter-rack links employ virtually imaged phased array (VIPA) technology, which enables flat, single-hop optical wireless paths across racks [16]. VIPA allows wavelength demultiplexing and precise beam steering, resulting in ultra-low-latency transmission across large-scale optical wireless deployments [16]. At the transmission control layer, we in-

corporate application-aware MAC protocols across both intra- and inter-rack DCN domains. These protocols dynamically allocate optical network resources based on diverse traffic priorities, ensuring that latency-sensitive applications, e.g., AI training, media streaming and AR/VR are served by the DCN preferentially achieving lower E2E latencies, while less time-critical workloads, e.g., cloud storage and background analytics are scheduled opportunistically.

Simulation-based evaluation of the proposed DCN architecture demonstrates near-optimal bandwidth efficiency, achieving up to 94% E2E utilization for server-to-server communication. Critically, latency-sensitive workloads (AI/ML) exhibit sub-us E2E latency, whereas less time-sensitive traffic (web services and cloud storage) remains below ms thresholds even under nominal load, satisfying typical application-level latency requirements. Owing to its flat, single-hop topology, the architecture exhibits strong scalability, supporting deployments exceeding 3000 servers without incurring significant additional control-plane complexity. Furthermore, server scaling experiments show a maximum latency increase of only 14%, indicating that the proposed DCN architecture can effectively accommodate growth in both computational resources and traffic demands. The remainder of this paper is organized as follows. In Section II, we present the proposed optical DCN architecture, detailing both inter-rack FSO and intra-rack fiber-based sub-networks. Section III provides a comprehensive performance evaluation based on simulation and scalability scenarios. Finally, Section IV concludes the paper and outlines directions for future research.

II. OPTICAL DCN ARCHITECTURE

The proposed DCN is organized into three complementary optical sub-networks, each optimized for a distinct DCN domain while ensuring seamless integration, as depicted in Fig. 1. The first, the inter-rack FSO sub-network, provides single-hop, wavelength-multiplexed links between ToR switches, enabling direct, low-latency communication between ToRs. The second, the intra-ToR sub-network, connects rack servers to their corresponding ToR via fiber, acting as aggregation of inter-rack traffic for forwarding through the inter-rack FSO network. The third, the intra-rack fiber-based sub-network, connects servers within the same rack, to serve intra-rack traffic.

A. VIPA-Based Inter-Rack FSO Sub-Network Topology

As depicted in Fig. 1, the inter-rack sub-network is responsible for all rack-to-rack communication i.e., inter-rack traffic, realized as a flat, single-hop interconnect fabric capable of supporting high data rates and low latencies across a large number of racks M . To this end, we leverage FSO links in conjunction with a VIPA device, which enables passive wavelength-to-space mapping and high-resolution angular dispersion [16].

Conventional FSO interconnects often rely on deterministic line-of-sight (LoS) paths or active beam steering mechanisms, even when passive components such as photonic integrated switches or metasurfaces are employed [9], [10]. These geometric constraints complicate deployment, scaling and main-

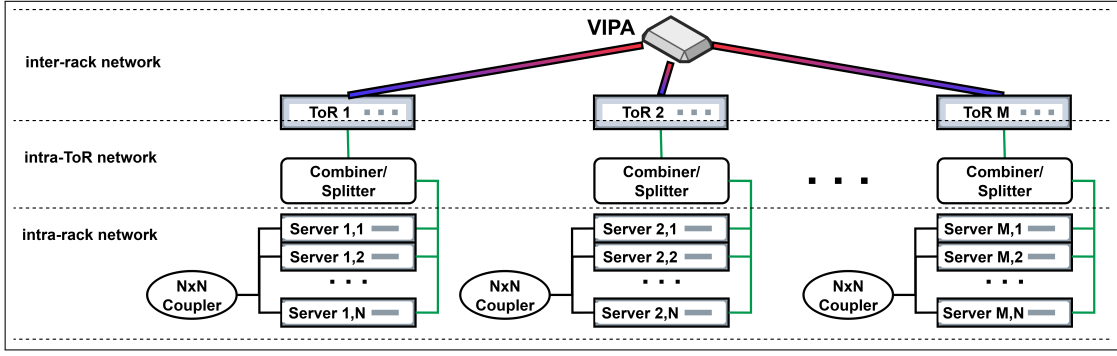


Fig. 1: Optical DCN architecture with M racks, each hosting N servers, comprising an inter-rack FSO sub-network (multicolor links), an intra-ToR server-ToR sub-network (green links) and an intra-rack server-to-server sub-network (black links).

tenance in dense DC environments. In contrast, the VIPA is an angularly dispersive optical component that effectively maps distinct optical wavelengths to unique spatial angles [16]. Specifically, when multiple wavelengths are input to the VIPA under controlled optical injection, interference among multiple internal reflections produces an output field in which each wavelength emerges at a resolvable angle in the far field [17]. This property enables scalability, as it allows each wavelength to be associated with a particular spatial direction and, by extension, a particular destination rack. By changing the transmission wavelength at a ToR, a corresponding beam is passively directed toward the intended rack without moving components or dynamic steering, substantially improving robustness and reducing control complexity. VIPA devices offer higher angular dispersion compared to arrayed waveguide gratings (AWGRs) or prism-based elements previously proposed for FSO-based DCNs [8], [15], enabling fine wavelength separation across many channels [18].

In our proposed DCN architecture, each rack contains a single ToR switch that interfaces with the inter-rack network through one optical transceiver. This transceiver includes a single wavelength-tunable laser source covering the inter-rack optical band, accompanied by a high-speed modulator and collimation optics. The tunable transmitters are capable of emitting on multiple wavelengths $\lambda_1^{\text{inter}}, \dots, \lambda_W^{\text{inter}}$, where W is the number of supported inter-rack channels. Collimating optics reduce beam divergence prior to injection into the VIPA, ensuring that the emergent beams remain sufficiently directed over typical DC distances of hundreds of meters [3]. After the VIPA, each wavelength travels along a distinct angular path through the data hall airspace; appropriately positioned and oriented receiving optics at each destination ToR collect the corresponding beam. The receiver front-end typically consists of a wide-aperture collection optic, aligned to the VIPA angular map and high-speed photodetectors [16]. By design, the ToR transceiver requires only a single physical interface to support all inter-rack wavelengths. In practice, the number W of resolvable channels (and therefore the corresponding maximum supported M) is determined by the optical spectrum, VIPA resolution and acceptable angular

separation to avoid inter-beam interference at the receiving end. Modern optical components support wavelength spacing on the order of 25–50 GHz across tens of nm of spectrum [19], [20], allowing practical deployment of 80–160 wavelength division multiplexing (WDM) inter-rack channels within the C-band. Theoretical studies indicate that, under idealized assumptions, high channel counts (up to approximately 1000) may be supported [16].

The proposed FSO interconnect is intended for short-range operation within DC halls, where free-space attenuation, atmospheric disturbances and channel impairments are relatively benign compared to long-haul outdoor FSO links [21]. The optical power budget of each inter-rack link accounts for transmitter optical output power, VIPA insertion loss, free-space path loss, beam divergence and receiver sensitivity. Collimating optics are chosen to balance beam divergence and receiver aperture size so that misalignment and thermal variations result in minimal degradation. While precise alignment of the VIPA output to receiving optics and a dust-free environment are required during installation, the passive nature of the VIPA and fixed mounting reduces the need for continuous active steering or alignment feedback loops.

B. Intra-ToR and Intra-Rack Sub-Network Topologies

Each rack comprises two dedicated optical sub-networks, inspired by the DCN architecture of [22]: (1) an intra-rack sub-network that supports direct server-to-server communication for intra-rack traffic and (2) an intra-ToR sub-network that aggregates inter-rack traffic from the rack servers to the ToR in the uplink (UL) direction and distributes traffic from the ToR to the rack servers in the downlink (DL) direction.

The intra-rack sub-network is based on a passive optical coupler interconnecting all N servers within a rack [22]. Each server is equipped with an intra-rack interface and connects to the coupler via an optical fiber carrying L data wavelengths and a dedicated control wavelength. Burst-mode transmitters and receivers, combined with semiconductor optical amplifiers (SOAs), enable low-loss, sub-ns switching, supporting high-bandwidth and low-latency server-to-server communication [22]. Inter-rack traffic generated by rack servers and destined to other racks is forwarded to the

ToR, which serves as the aggregation point and gateway before traversing the inter-rack sub-network. To this end, each server is equipped with a separate intra-ToR interface employing fixed-wavelength transceivers, complemented by SOAs to support ns-scale switching. The ToR switch hosts a corresponding intra-ToR interface with fixed-wavelength transceivers.

C. Application-Aware MAC Scheduling Protocols

Based on the traffic taxonomy introduced in [2] and the application trends discussed in Section I, modern DCN workloads can be broadly classified into four representative application classes: (1) AI/ML workloads for training and inference (AI/ML), (2) immersive media services (AR/VR/XR), (3) cloud storage and backend services (Storage) and (4) web-based communications (Web APIs). Among these, AI/ML traffic is the most latency-sensitive, with typical E2E latency requirements ranging from 0.02 to 5 ms [23]. AR/VR/XR applications also impose stringent latency constraints, generally between 1 and 20 ms [23]. Storage workloads exhibit higher tolerance, with latency bounds between 5 and 30 ms, while Web API traffic can tolerate delays on the order of several tens of ms [23]. To accommodate this traffic heterogeneity across the DCN, we adopt an application-aware MAC scheduling strategy for each sub-network that explicitly prioritizes latency-sensitive traffic classes. Transmissions are organized into recurring control and data cycles. During a control cycle, control information is exchanged and scheduling decisions are computed ahead of data transmission, which is performed in the upcoming data cycles.

Inter-rack traffic is generated by rack servers, stored into the server's four inter-rack buffers (one for each traffic application class) and is transmitted to the corresponding ToR via the intra-ToR sub-network. At the ToR, it is aggregated to the ToR's four application buffers and is subsequently transmitted to the destination rack via the VIPA-based FSO fabric. A centralized Software-Defined Networking (SDN) controller coordinates inter-rack communication by collecting transmission requests from all ToRs over the control wavelength during each control cycle. Each request includes the source and destination ToR identifiers, packet size and priority class. The SDN controller applies a bandwidth-opportunistic MAC scheduling algorithm that accounts for the physical constraints of the FSO links (i.e., a single inter-rack transmitter and receiver per ToR - see Section II-A). Within these constraints, higher-priority traffic with larger payloads is scheduled first to maximize bandwidth utilization and enable QoS provisioning, followed by lower-priority traffic. The computed schedule is subsequently distributed to the ToRs over the control wavelength during the control cycle and is subsequently executed during the next data cycle, ensuring contention-free, high-throughput inter-rack communication.

Intra-rack traffic on the other hand is stored in four dedicated server buffers, one per application class and is then handled using a broadcast-and-select WDM MAC protocol inspired by [24] over L intra-rack channels and a dedicated

intra-rack control wavelength. Transmissions occur in synchronized cycles separated by guard bands to accommodate laser turn-on/off times. Control information (source and destination server identifiers, packet size, and priority class) is broadcasted by all servers over the shared intra-rack control wavelength. This enables each server to independently compute its local transmission schedule for the upcoming data cycles with the following policy. In each control cycle a traffic priority matrix is generated that determines the transmission order for the subsequent data cycles over the L intra-rack wavelengths. Servers are randomly hierarchized (for fairness) and thereby schedule higher-priority packet transmission over lower-priority ones. If a scheduled server does not have sufficient buffered traffic to fully utilize its allocated transmission opportunity, the unused capacity is reassigned to the next eligible server in the priority matrix, with additional protocol details available in [24].

Finally, the intra-ToR sub-network employs a distinct MAC protocol for UL and DL communication, where two wavelengths are used (UL, DL). In the UL direction, a time-division multiplexing (TDM)-based scheduling scheme is used [22]. Servers advertise pending inter-rack transmission requests via control packets over the shared intra-rack control channel, similarly to the intra-rack case. Based on this information, servers compute their UL transmission schedules in a distributed manner, prioritizing latency-sensitive traffic classes. In the DL direction, transmissions from the ToR to servers follow a comparable cyclic operation, with the ToR acting as the sole transmitter over fixed-wavelength channels in a one-to-many fashion.

III. PERFORMANCE EVALUATION

This section evaluates the proposed optical DCN via simulation methodology, key parameters and metrics, followed by baseline performance, scalability and SotA comparisons.

A. Evaluation Methodology

To evaluate our architecture, we developed a Python-based simulator that models all communications within the DCN [22]. Traffic is generated at the server level using Trafpy [25], with 80% of traffic remaining intra-rack and 20% inter-rack [24]. Four coexisting application classes are considered (see Section II-C), ordered from highest to lowest priority: (1) AI/ML, (2) AR/VR/XR, (3) Storage and (4) Web APIs. The corresponding traffic mix is set to 40%, 30%, 20%, and 10%, respectively, based on the insights from [2]. The baseline topology consists of $M = 16$ racks, each equipped with $N = 16$ servers, for a total of 256 servers. Server and ToR buffers (incurring negligible cost) are fixed at 1 MB. Propagation delays are set to 16.7 ns for the intra-rack and 500 ns for the inter-rack communication (corresponding to physical distances of 2.5 m between rack servers and 75 m between racks [22]). Four data wavelengths are employed for intra-rack communication ($L = 4$), while two wavelengths (UL, DL) are used in the intra-ToR sub-network. In the inter-rack sub-network, the number of wavelengths equals the number of racks $W = M$ in each simulation scenario.

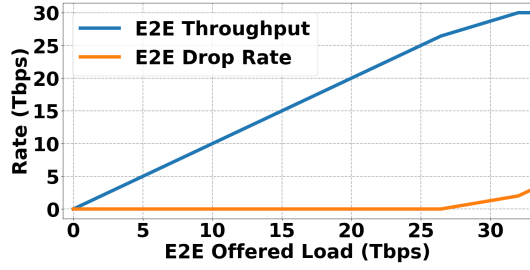


Fig. 2: Baseline topology (16×16): E2E throughput & drop rate as a function of total offered DCN load

All data and control links operate at 400 Gbps, yielding an aggregate E2E capacity of 32 Tbps. Guard bands of 90 bytes are assumed, based on a SOA switching time of 900 ps [24]. Performance is evaluated in terms of throughput, packet drop rate, throughput efficiency and E2E latency.

B. Simulation Results

Figure 2 shows the E2E throughput and packet drop rate as a function of offered load across the DC. The results demonstrate that the proposed DCN solution achieves near-optimal throughput across all sub-networks. Throughput efficiency remains at 100% for offered loads up to 26.4 Tbps, saturating at higher loads with a maximum efficiency of 94% at 32 Tbps. Packet drop rate remains negligible ($\sim 0\%$) until the system approaches full saturation, rising modestly to 6% at the nominal offered load. This behavior highlights the seamless integration of intra-rack, intra-ToR and inter-rack sub-networks, ensuring that no bottlenecks occur across the three sub-networks and that high-volume traffic is efficiently routed across racks without congestion.

Figure 3 reports the E2E latency per application vs. normalized offered load. At 25% load, latency-critical AI/ML traffic experiences an average E2E latency of $0.3 \mu\text{s}$, increasing only to $0.4 \mu\text{s}$ at 100% load; a modest 33% rise (well below the 0.02–5 ms target specified in Sec. II-C), confirming that high-priority workloads remain effectively isolated from congestion even under full saturation. AR/VR/XR traffic grows from $4.7 \mu\text{s}$ at 25% load to $9.1 \mu\text{s}$ at the nominal load, representing a 93% increase; this remains comfortably below the 1–20 ms latency range defined for interactive media workloads (Sec. II-C). Storage traffic increases from $5.78 \mu\text{s}$ to $362.8 \mu\text{s}$ (6200%), still well below the 5–30 ms expected range, demonstrating that lower-priority but bandwidth-intensive workloads are delayed in an acceptable manner. Web API traffic exhibits a rise from $6.91 \mu\text{s}$ to $493.6 \mu\text{s}$ (7000%), remaining safely within the tens-of-ms tolerance specified for this class (Sec. II-C). These trends illustrate deliberate latency differentiation: critical workloads maintain ultra-low, near-constant latency, while lower-priority traffic is incrementally delayed under high load without violating acceptable bounds.

Figure 4 shows how system scaling affects E2E latency across application classes under nominal load. As the number of racks (M) increases, the number of inter-rack wavelengths

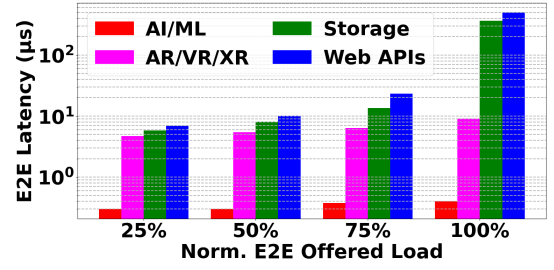


Fig. 3: Baseline topology (16×16): E2E latency per application as a function of normalized offered DCN load

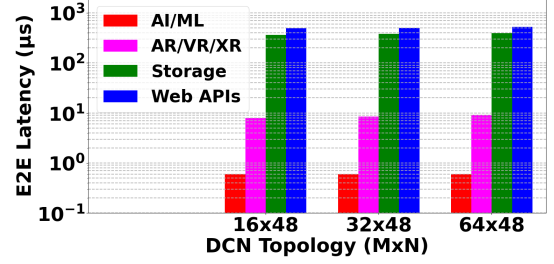


Fig. 4: E2E latency per application as a function of DCN topology scale (M racks, N servers per rack)

(W) must also increase, since our design sets $W = M$. For the 16×48 topology, AI/ML traffic exhibits an E2E latency of $0.6 \mu\text{s}$, AR/VR/XR $8.0 \mu\text{s}$, Storage $365.4 \mu\text{s}$ and Web APIs $494.1 \mu\text{s}$, all well within the latency bounds defined in Sec. II-C. As system increases to 32×48 , latency increases marginally: AI/ML remains at $0.6 \mu\text{s}$ (0%), AR/VR/XR rises to $8.5 \mu\text{s}$ (+6.3%), Storage to $378.6 \mu\text{s}$ (+3.6%) and Web APIs to $502.2 \mu\text{s}$ (+1.7%). Further scaling to $M \times N = 64 \times 48$ (i.e., over 3K servers) results in slight increases: AI/ML $0.6 \mu\text{s}$ (0%), AR/VR/XR $9.1 \mu\text{s}$ (+13.8%), Storage $393.9 \mu\text{s}$ (+7.8%) and Web APIs $522.7 \mu\text{s}$ (+5.8%). These limited increases highlight that even for large-scale deployments, all application classes remain comfortably within their respective service-level expectations. Robust scalability is enabled by cross-layer traffic coordination across intra- and inter-rack domains, and priority-based MAC protocols that dynamically allocate bandwidth while ensuring QoS and high utilization. Additionally, using a single transceiver per ToR (rather than four as in our previous work [22]) reduces inter-rack energy consumption by 75%, assuming all other parameters are equal.

To further evaluate the proposed DCN, we simulate a setting with $M = 54$ racks and $N = 48$ servers per rack, matching the scale of the recent FSO-based SotA architecture ODRAD [15] (see also Section I). Table I summarizes the comparison in terms of throughput efficiency and E2E latency across different application classes at nominal offered load. Our results indicate that our proposed architecture achieves a throughput efficiency of 93%, significantly higher than the 78% reported for ODRAD. This improvement can be attributed to the seamless integration of the fiber-based intra-rack and VIPA-enabled inter-rack sub-networks along with

TABLE I: Performance Comparison of the Proposed DCN and ODRAD Architectures

Metric	Application	Proposed DCN	ODRAD [15]
Servers per Rack		48	48
Total Servers		2592	2560
Throughput Efficiency (%)		93	78
	AI/ML	0.6	7–10
E2E Latency for	AR/VR/XR	9.05	7–10
Max Load (μ s)	Storage	392.4	7–10
	Web APIs	520.1	7–10

the proposed MAC protocols, which collectively eliminate bottlenecks in both local and rack-to-rack communication, as previously highlighted in Figures 2 and 4. In terms of E2E latency, our application-aware MAC scheduling ensures that time-critical workloads, particularly AI/ML traffic, experience ultra-low latency (0.6 μ s), orders of magnitude lower than ODRAD’s 7–10 μ s range. Media-related traffic (AR/VR/XR) also benefits from prioritized treatment, maintaining competitive latency comparable to ODRAD. Conversely, less time-sensitive workloads such as Storage and Web APIs experience higher latencies (392.4 μ s and 520.1 μ s, respectively), reflecting the deliberate design choice to favor latency-critical applications. This selective prioritization contrasts with ODRAD’s uniform treatment across all traffic classes, demonstrating that the proposed DCN solution not only maximizes throughput but also promotes sustainable, performance-aware operation, ensuring that high-priority workloads meet stringent latency requirements while less critical traffic is gracefully handled under nominal load conditions.

IV. CONCLUSION

In our work, we have presented a scalable hybrid optical DCN architecture that integrates intra-rack fiber-optical links with inter-rack VIPA-based free-space optical connections, combined with application-aware MAC scheduling. Simulation results demonstrate near-optimal bandwidth efficiency and strict latency differentiation, achieving sub-us end-to-end latency for time-critical DC workloads while maintaining acceptable performance for less time-sensitive traffic, thereby illustrating the potential of application-aware optical fabrics for next-generation DCNs. Future research includes system scaling via layered VIPA-based interconnects, exhaustive evaluation with real DC traffic traces, comparison with SotA DC architectures, including leaf-spine and Optical Circuit Switching (OCS)-based solutions and development of a proof-of-concept to validate the proposed architecture.

REFERENCES

- [1] The networks that will turn potential into profit in 2026. [Online]. Available: <https://www.computerweekly.com/news/366636882/The-networks-that-will-turn-potential-into-profit-in-2026>
- [2] P. A. Baziana, “Optical data center networking: A comprehensive review on traffic, switching, bandwidth allocation, and challenges,” *IEEE Access*, vol. 12, pp. 186 413–186 444, 2024.
- [3] Data centre capacity set to double, JLL global data center outlook 2026. [Online]. Available: <https://capacityglobal.com/news/data-centre-capacity-set-to-double-jll/>
- [4] Why fiber optics is replacing copper in data centers. [Online]. Available: <https://www.eetimes.com/why-fiber-optics-is-replacing-copper-in-data-centers/>
- [5] The trends shaping the data center in 2026 and beyond. [Online]. Available: <https://www.siemon.com/en/the-trends-shaping-the-data-center-in-2026-and-beyond/>
- [6] K.-i. Sato, “Optical switching will innovate intra data center networks [invited tutorial],” *Journal of Optical Communications and Networking*, vol. 16, no. 1, pp. A1–A23, 2023.
- [7] Z. Xie and Q.-F. Yang, “Terabits without fibres,” *Light: Science & Applications*, vol. 14, no. 1, p. 232, 2025.
- [8] Y. Yi *et al.*, “Increasing capacity of intra-datacenter communications using a novel combination of DQPSK and PAM-4 modulation formats,” *Optical Fiber Technology*, vol. 94, p. 104309, 2025.
- [9] S. Zhang *et al.*, “Photonic integrated multicast switch-based optical wireless data center network,” *Journal of Optical Communications and Networking*, vol. 15, no. 7, pp. C54–C62, 2023.
- [10] W. Qiu *et al.*, “Universal strategy study of applying passive metasurfaces to an intra-dc wireless-optical interconnection,” *Journal of Optical Communications and Networking*, vol. 16, no. 10, pp. 929–940, 2024.
- [11] H. Kazemi *et al.*, “Terabit optical wireless link design for inter-rack communication in 6G data center networks,” in *2023 IEEE International Conference on Communications Workshops (ICC Workshops)*. IEEE, 2023, pp. 1771–1776.
- [12] S. Zeng *et al.*, “Solid-state 360° optical beamforming for reconfigurable multicast optical wireless communications,” *Optics Express*, vol. 31, no. 6, pp. 10 070–10 081, 2023.
- [13] T. Lin *et al.*, “OWC4DT: a 25 Gbps plug-and-play self-alignment optical wireless communication system for distributed training in data center,” in *ICC 2025-IEEE International Conference on Communications*. IEEE, 2025, pp. 3107–3112.
- [14] Q. Kong *et al.*, “Scalable matrix ring data center interconnect based on FSO technology,” *Optics Express*, vol. 33, no. 12, pp. 25 016–25 038, 2025.
- [15] K. Geresu *et al.*, “ODRAD: An optical wireless DCN dynamic-bandwidth reconfiguration with AWGR and deep reinforcement learning,” *Optical Switching and Networking*, vol. 52, p. 100771, 2024.
- [16] S. Xiao and A. M. Weiner, “2-D wavelength demultiplexer with potential for 1000 channels in the C-band,” *Optics Express*, vol. 12, no. 13, pp. 2895–2902, 2004.
- [17] C. Cheng *et al.*, “Net 1 Tbps multi-user indoor FSO downlink over 25 m based on cost-effective point-to-multipoint coherent optics and probabilistic shaping,” in *2023 Optical Fiber Communications Conference and Exhibition (OFC)*. IEEE, 2023, pp. 1–3.
- [18] X. Zhu *et al.*, “Dispersion characteristics of the multi-mode fiber-fed VIPA spectrograph,” *The Astronomical Journal*, vol. 165, no. 6, p. 228, 2023.
- [19] DWDM channel spacing and ITU-T G.694.1 frequency grid. [Online]. Available: https://mapyourtech.com/wp-content/uploads/resources_guides/infographics_wdm.html
- [20] F. El-Nahal *et al.*, “A bidirectional wavelength division multiplexed (WDM) free space optical communication (FSO) system for deployment in data center networks (DCNs),” *Sensors*, vol. 22, no. 24, p. 9703, 2022.
- [21] S. Phuchortham and H. Sabit, “A survey on free-space optical communication with rf backup: Models, simulations, experience, machine learning, challenges and future directions,” *Sensors*, vol. 25, no. 11, p. 3310, 2025.
- [22] G. Drainakis *et al.*, “Scalable and low server-to-server latency data center network architecture based on optical packet inter-rack and intra-rack switching,” *Journal of Optical Communications and Networking*, vol. 15, no. 11, pp. 804–819, 2023.
- [23] Hexa-X-II. (2024) 6G use cases and requirements: HEXA-X-II d1.2 deliverable. [Online]. Available: https://hexa-x-ii.eu/wp-content/uploads/2024/01/Hexa-X-II-D1.2_slideSet.pdf
- [24] G. Drainakis *et al.*, “Application-aware resource allocation and traffic classification for us-latency optical data centers,” in *2025 IEEE International Mediterranean Conference on Communications and Networking (MeditCom)*. IEEE, 2025, pp. 1–6.
- [25] C. W. Parsonson *et al.*, “Traffic generation for benchmarking data centre networks,” *Optical Switching and Networking*, vol. 46, p. 100695, 2022.