

一种基于图表示学习的光网络算力调度方法

于添阔,姚秋彦*,杨辉,龚盛业

(北京邮电大学 信息光子学与光通信全国重点实验室,北京 100876)

摘要:针对现有算力资源与光网络独立管理导致的资源调度不协调问题,提出了一种基于图表示学习的光网络算力调度方法。该方法通过构建节点和边特征图并利用图卷积网络进行聚类,形成二分图来映射算力业务的源节点到目的节点。通过引入学习因子优化二分图中的映射关系,以最小化时延为目标,实现高效的算力调度。仿真结果表明,所提方法在多粒度算力业务共存环境下显著地降低了阻塞率和业务平均处理时延,提高了资源利用率。

关键词:弹性光网络;算力调度;图表示学习

中图分类号:TN915 文献标志码:A 文章编号:1002-5561(2024)05-0046-05

DOI:10.13921/j.cnki.issn1002-5561.2024.05.008

开放科学(资源服务)标识码(OSID):



Optical network computing power scheduling method based on graph representation learning

YU Tiankuo, YAO Qiuyan*, YANG Hui, GONG Shengye

(State Key Laboratory of Information Photonics and Optical Communication,
Beijing University of Posts and Telecommunications, Beijing 100876, China)

Abstract: A graph representation learning based on optical network computing power scheduling method is proposed to address the problem of resource scheduling mismatch caused by the independent management of existing computing power resources and optical networks. This method constructs node and edge feature maps and uses graph convolutional networks for clustering, forming a bipartite graph to map the source nodes of computing power services to the destination nodes. By introducing learning factors to optimize the mapping relationship in the bipartite graph, with the goal of minimizing latency, efficient computing power scheduling can be achieved. The simulation results show that the proposed method significantly reduces the blocking rate and average processing delay of services in a multi granularity computing power coexistence environment, and improves resource utilization.

Key words: elastic optical network, computing power scheduling, graph representation learning

0 引言

随着第五/第六代(5G/6G)移动通信技术的快速发展和新兴算力应用的不断涌现,通信网络正经历着从仅提供传送管道到同时提供信息传输与信息处理能力的转变,以光网络作为承载底座的大容量高可靠

算力网络应运而生^[1]。然而,算力与光网络在管控上相互隔离,且多样化的算力业务带宽跨度大,呈多颗粒度形态,经粗大粒度的光传送节点聚合后会引入额外的时延,严重影响了差异化算力业务的服务质量^[2-3]。因此,如何利用光网络实现高效的算力调度成为制约算力网络发展的关键瓶颈。弹性光网络凭借其频谱资源的细粒度划分及灵活可调的特点,可以为多粒度的算力业务提供更加精细的传输配置^[4-5]。由于异构的热点流量在时域上的差异,城域光网络的业务流量具有突发性,经过光传输节点汇聚后,容易引发阻塞进而增加额外的时延^[6]。此外,由于多优先级业务的目的节点不一致,在路由和频谱资源分配阶段会引入不同的时延约束。特别是在分配过程中,相邻节点的业务流

收稿日期:2024-02-28。

基金项目:国家自然科学基金项目(62201088)资助。

作者简介:于添阔(1999—),男,博士研究生,现就读于北京邮电大学电子工程学院电子科学技术专业,主要研究方向为弹性光网络、算力网络、深度学习。

*通信作者:姚秋彦(1989—),女,博士,特聘副研究员,主要研究方向为弹性光网络 and 智能光网络。



向呈现出复杂的相关聚合特征,加之算力节点资源状态快速变化,导致了响应时间的不确定性。因此,如何利用好算力光网络这种天然的图数据结构,进行低时延的调度显得尤为重要。

为了解决上述问题,本文提出一种基于图表示学习的光网络算力调度(GRL-ONCPS)方法。

1 基于 GRL 的调度架构和弹性光网络定义

光网络的复杂拓扑结构和动态业务特性使得传统调度方法难以有效适应其需求。光网络中的节点和边可以被视为网络拓扑中的实体和连接,形成了一个复杂的图形结构。而图卷积神经网络(GCN)具备捕捉复杂图结构中关系的能力,能够学习节点之间的相互影响和整个网络的全局特征^[7-8]。除此之外,GCN 具备很好的处理动态、未知目标节点的能力,可以灵活适配光网络中动态变化的业务需求。为了最大限度地利用算力资源,同时匹配光网络的高动态性,本文在传统光传送节点的基础上引入算力智能板,突破算力和网络管控隔离的僵化配置制约,实现传输时的实时流量检测和带宽动态可调。GRL-ONCPS 系统架构图如图 1 所示。

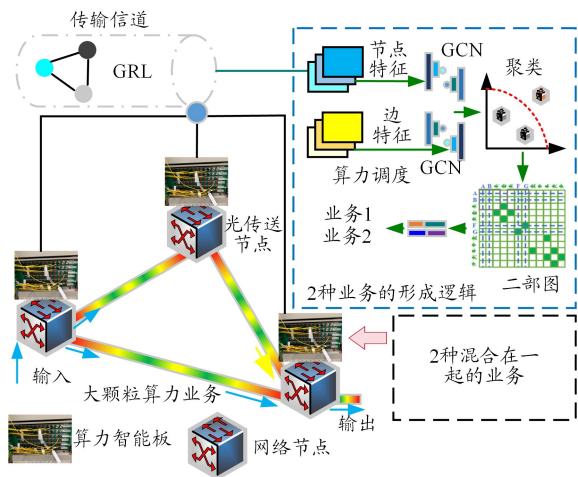


图 1 GRL-ONCPS 系统架构图

该架构在交换机空闲卡槽中嵌入了算力智能板,旨在赋予系统高性能动态计算能力,以精准满足智能调度的复杂需求。在业务传输过程中,算力智能板首先通过光网络中的算力设备读取数据,随后采用基于 GCN 模型适配的 GRL-ONCPS 方法对数据进行聚类处理。紧接着,学习模块对算力调度进行映射和优化决策,算力智能板则根据这些决策实时处理并下发指令,以动态调整流经光传送节点的大颗粒度算力业务流量。最后,控制层利用南向接口及 OpenFlow 协议采

集底层数据流量信息,并维护检测网络状态信息矩阵,进而执行路由和频谱资源分配算法的计算,随后将计算结果下发至算力智能板及光交换模块,以执行相应的算法调控。因此,通过智能板模块的实时监测,可以更精确地感知网络状态,实现对算力的动态调度,从而更好地适应光网络中高动态性的数据传输需求,实现网络资源的最优利用,提升了网络的整体性能和服务质量。

GRL-ONCPS 方法的输入涵盖了丰富的网络状态信息,包括实时流量变化、相邻节点流量动态、不同粒度数据包的业务特征、节点间连接关系及链路带宽等。在输出端,该方法则通过数据包的智能交换与重组,实现弹性光网络中的重路由优化与频谱碎片重构。

本文将弹性光网络表示为 $G(V, E, N, C)$ 。其中, V 表示节点集合, $V_k = \{V_k^a, V_k^p\}, k=1, 2, \dots, M$, 即共有 M 个算力节点和光传送节点, a 表示传输节点, p 表示计算处理节点; E 表示链路集合, N 表示每根光纤中的频谱槽个数; C 表示算力资源, $C_k = \{C_k^c, C_k^m\}$, m 为显存容量, c 表示计算资源,即节点算力资源包括计算资源和存储资源。此外,将算力业务表示为 $R_i\{S_i, D_i^S, B_i, r_e, r_t, t_s\}$ 。其中, S_i 是业务源节点, $S_i \in V$; D_i^S 是为业务 i 提供算力服务的目的节点, $D_i^S \in V$, j 为目的的节点数, $j \in [M]$, $S_i \subseteq [M]$, $[M]$ 被定义为 $1 \sim M$ 的自然数集合, B_i 为业务 i 需要的带宽; r_e 为业务 i 的算力需求量, r_t 为业务 i 的时延容忍度, t_s 为业务 i 的到达时间。算力调度就是为业务 R_i 在当前的网络状态 G_i 找到目的节点 D_i^S , 占用算力资源 C_k^m 。系统进行计算卸载和状态更新,得出新的网络状态 G^* 。

2 GRL-ONCPS 方法原理

本文提出的 GRL-ONCPS 方法主要分为 2 个阶段。首先,使用 GCN 按照资源占用程度对节点进行聚类划分;其次,通过二部图模型来进行算力提供方选择、路由以及频谱资源分配,从而为每一个算力业务需求提供最终的算力调度方案。

2.1 基于 GCN 的资源聚类

首先,构建衡量资源占用程度的指标,包括算力资源、频谱资源占用度和最大空白频谱槽。其中,算力资源是在点层次上的信息,频谱资源是在边层次上的信息。在图 1 中,本文建立 2 张辅助图进行学习:一张

于添阔,姚秋彦,杨辉,等:一种基于图表示学习的光网络算力调度方法

是原网络的拓扑,将边信息转化为点信息,同时输入到特征向量 X_1 中,并将邻接矩阵 A^1 和特征向量 X_1 输入到 GCN 中对点进行二次聚类,得出光网络中节点的聚类结果 $\{U_1, T_1\}$ 。另一张拓扑是针对边信息进行抽象,在光网络拓扑中以边之间是否有且只有一个点为判据,将边转换为点,从而重构出新的拓扑。接着,将点信息再次转化为边信息,并将其输入到特征向量 X_2 中。随后,将邻接矩阵 A^2 和特征向量 X_2 输入到 GCN 中,对点进行二次聚类,得出光网络中链路的聚类结果 $\{U_2, T_2\}$ 。其中, $U_1=P, D, Q, H, J, K, L$ (节点名称), $T_1=F, B, A, G, I$ (节点名称), $\{U_1, T_1\}$ 表示光网络中点的簇划分, U_1 为资源充足的簇, T_1 为资源不充足的簇。 $U_2=6, 2, 8, 1, 11, 3, 9, 7, 12, T_2=4, 5, 10, 14, 13$, 而 $\{U_2, T_2\}$ 表示光网络中链路的簇划分, U_2 为资源充足的簇, T_2 为资源不充足的簇。

2.2 基于二部图的资源映射与调度

本文得到光网络点和边的簇划分 $\{U_1, T_1\}$ 和 $\{U_2, T_2\}$ 后,这里的算力调度问题可以视为寻找从 T_1 到 U_1 的多对多映射关系,如表 1 所示。 $T_1 \rightarrow U_1$ 中下标 1 表示 2 个点之间有箭头。不失一般性,将其记作 $W_1 = [T_1 \rightarrow U_1]$ 。

表 1 $T_1 \rightarrow U_1$ 二部图映射表

U_1	T_1			
	T_1^1	T_1^2	T_1^3	$T_1^{\max}$
U_1^1	1	0	1	1
U_1^2	0	1	0	1
U_1^3	1	0	1	0
$U_1^{\max}$	1	1	1	0

注:在 T_1 的 16 个数值中,1 表示有算力调度,0 表示无算力调度。

设置学习因子 W 和 W_1 同维,且 $W \in (-1, 1)$, $W_1 = u(W)$, $u(\cdot)$ 为阶跃函数。

此外,光网络链路上的距离代价矩阵可表示为 $\Phi_{M \times M}$,而由 GCN 对边拓扑聚类得出的边上的簇划分为频谱资源的紧张程度信息,它可以用来更新选路代价。本文将簇划分之前的特征向量输入到一个全连接层,进而使用 Softmax 激活函数将其量化为 0~1,并按照边索引在原光网络拓扑上进行扩展,由此可以更新选路的代价矩阵 $Y_{p,q}$, p, q 分别为 $Y_{p,q}$ 的行数和列数。

$$Y_{p,q} = \Phi_{p,q} L_{p,q} \quad (1)$$

其中, $\Phi_{p,q}$ 、 $L_{p,q}$ 分别表示距离代价矩阵扩展后的更新张量。这时需要由 $p=S_i$ 到 $q=D_i^S$ 依据 Dijkstra 算法寻找

最短距离,得出最佳传播路径 P_i 和最小传播时延 $\tau_i^{d,j}$,再对该业务的所有路径求和可以得出传播总时延 τ_i^d 。

$$\tau_i^d = \sum_j \tau_i^{d,j} \quad (2)$$

然而,除了传播时延之外,该业务的处理时延 τ_i 还包括计算时延 τ_i^c 。

$$\tau_i = \tau_i^d + \tau_i^c \quad (3)$$

这里的 τ_i^c 可以被表示为

$$\tau_i^c = \sum_j \frac{r_{e_{k=j}}}{C_{k=j}} \quad (4)$$

最后,设置损失函数 ε 为

$$\varepsilon = \frac{\sum_i ||r_i - \tau_i||^2}{n} + l(r_i < \tau_i) \$ \quad (5)$$

其中, $||\cdot||^2$ 为 L2 范数, n 为该时刻集合中的业务数, $l(\cdot)$ 为示性函数, $\$$ 为惩罚项, τ_i 为业务的处理时延。

3 仿真结果分析

3.1 仿真设置

为了验证所提出的 GRL-ONCPS 方法在算力调度任务上的性能,本文在算力调度场景下生成业务进行实验,将该方法与传统的均匀算力目标节点首次命中(Homogeneous-FF)方法进行对比分析,得出了在不同负载和任务量下的阻塞率 β 、平均时延 $\bar{\tau}$ 、资源利用率 ρ 、服务质量 Q 等性能的分析结果,如式(6)~式(9)所示。

$$\beta = \Pi / Z \quad (6)$$

其中, Z 为业务总数, Π 表示阻塞的业务数。

$$\bar{\tau} = \frac{\sum_i \tau_i}{Z} \quad (7)$$

$$\rho = \frac{\sum_{q,l} F_l^q}{N \sum_{R_u \in S} e_u} \quad (8)$$

其中, e_u 为业务路径的总边数, S 为被阻塞业务的集合, F_l^q 为资源占用矩阵。

$$Q = \frac{\beta^{\bar{\tau}} \times \sum_i \check{r}_{\tau_i}}{Z} \quad (9)$$

其中, $\check{r}_i$ 为满足时延容忍度的业务, t_i 为业务 i 服务完成时刻。

本文提出的算力调度场景为城域弹性光网络,采用 66 个节点、70 条边的拓扑进行仿真^[9]。设置每条边的最大频谱槽数为 100 个,采用自适应距离调制方式。

$$N_i^R = \frac{B_i}{12.5\varphi} \quad (10)$$

其中, N_i^R 为 R_i 占用的频谱槽个数, φ 为调制阶数,可在 1~5 内变化。例如,对于正交幅度调制(32QAM)方式, $\varphi=5$ 。因此,这种设置可以模拟多种类型的算力带宽业务,覆盖范围从 12.5~625 Gb/s。通过这种方式,本文验证了算法在城域网络场景下的可行性。

随机生成来自不同源节点的业务 ($M=5\ 000$), 每项业务所需的频谱槽在 1~10 内随机产生, 并采用距离自适应调制方式。所需的算力资源也在 1~5 个单位之间随机产生, 而业务的时延容忍度则在 1~6 内随机产生。各节点的算力资源容量分别为 {1 000, 1 200, 1 352, 2 451, 3 612, 4 751, 6 821, 5 410, 1 025, 4 100, 1 254, 3 541, 2 451, 4 751, 6 821, 5 410, 1 025, 4 100, 1 254, 3 541, 2 451, 1 000, 1 200, 1 111, 1 352, 2 451, 3 612, 4 751, 6 821, 5 410, 1 025, 4 100, 1 254, 3 541, 2 451, 4 751, 6 821, 5 410, 1 025, 4 100, 1 254, 3 541, 2 451, 1 000, 1 200, 1 111, 1 352, 2 451, 3 612, 4 751, 6 821, 5 410, 1 025, 4 100, 1 254, 3 541, 2 451} 千个单位, 处理能力分别为 {2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 1, 4, 5, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9, 2, 5, 4, 6, 5, 1, 2, 3, 7, 8, 9} 个单位/ms。此外, 各节点的流量负载从 200~1 200 Erlang 进行遍历, 业务经过节点转发的消耗时延为 0.5 ms。根据上述场景参数设定, 本文的仿真测试是在配备有第 11 代 Intel® Core™ i7-11800H @ 2.30 GHz 处理器和 16 GB RAM 的 Anaconda 虚拟机环境下使用 PyTorch 框架进行的。

3.2 结果分析

当流量负载为 200 Erlang 时, 不同方法的阻塞率情况如图 2 所示。可以看出, GRL-ONCPS 方法的阻塞率明显低于传统的 Homogeneous-FF 方法。随着业务数量的增加, GRL-ONCPS 方法的性能表现得越来越好, 这主要是因为它能够很好地学习到图结构中相邻节点的状态相关性, 并对特定路由方向保留较大的剩余空白频谱槽, 从而在差异化的多颗粒业务共存时能有效地避免阻塞。

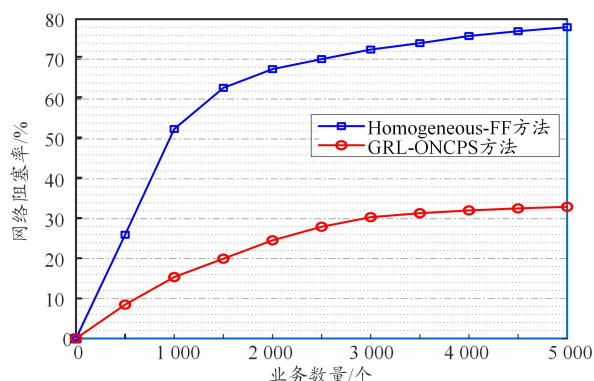


图2 不同方法的阻塞率

图3为流量负载在200 Erlang时的算力和频谱资源利用率的对比曲线。可以看出,随着业务数量的增加, GRL-ONCPS方法的资源利用率明显高于 Homogeneous-FF方法。这主要是因为 GRL-ONCPS方法采用了分布式部署的算力智能板, 在业务动态流经弹性光网络传送节点时进行带宽整合和频谱重构, 从而提高了整体的算网资源利用率。

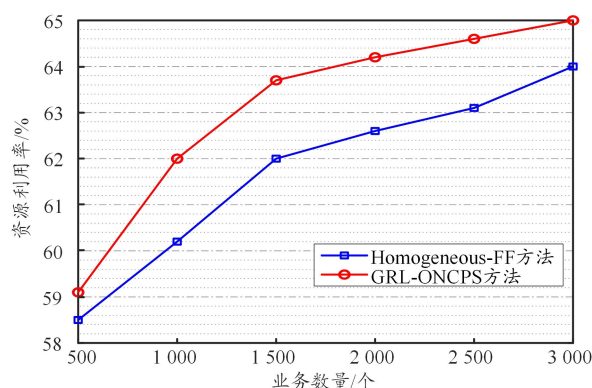


图3 不同方法的资源利用率

当网络负载为200 Erlang时, 2种方法的服务质量和平均处理时延如图4所示。可以看出, GRL-ONCPS方法显著提高了用户的服务质量, 并降低了平均处理时延。为了验证模型的鲁棒性, 本文改变了网络负载, 并得到了不同方法的阻塞率情况, 如图5所示。可

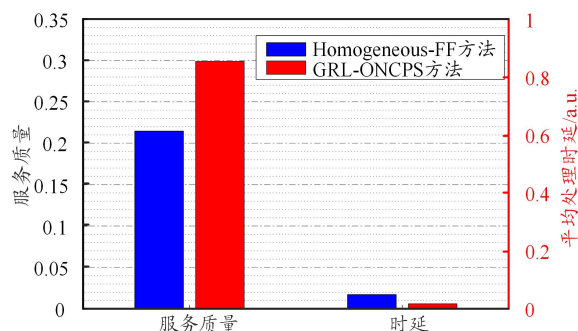


图4 不同方法的服务质量和平均时延

于添阔,姚秋彦,杨辉,等:一种基于图表示学习的光网络算力调度方法

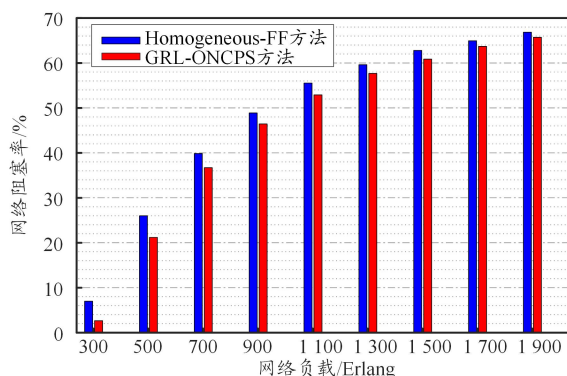
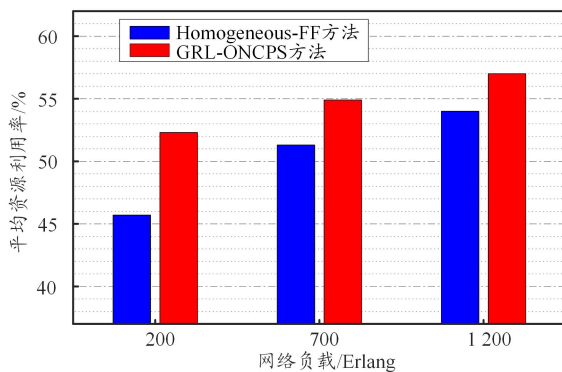


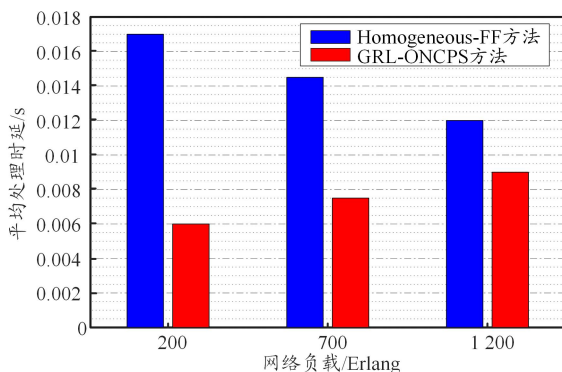
图5 不同负载下各方法的阻塞率

以看出,随着负载的增加,2种方法的阻塞率整体呈现上升趋势,这与实际情况相符;同时,随着负载的增加,网络状态趋于饱和。而GRL-ONCPS方法的阻塞率低于Homogeneous-FF方法,这表明GRL-ONCPS方法具有更好的性能。

为了更加清晰直观地展现不同负载下的性能,本文仿真对比了2种方法在不同负载下的平均资源利用率变化情况,如图6(a)所示。可以看出,2种方法的平均资源利用率随着负载的增加而增加,且GRL-ONCPS方法的平均资源利用率更高。图6(b)为上述不同



(a) 平均资源利用率



(b) 平均服务时延

图6 性能对比

方法在平均处理时延方面的对比情况。可以看出,与Homogeneous-FF方法相比,GRL-ONCPS方法能够显著降低平均处理时延,为多颗粒算力业务在光网络中的算力调度提供了低时延的快速承载基础。

4 结束语

本文提出了GRL-ONCPS方法,在典型的域域光网络场景下进行了仿真验证,仿真结果表明该方法能够显著降低阻塞率和平均处理时延,并且提高了光网络中的算力和频谱资源利用率。

鉴于大多数现实网络都是动态变化的,而算力业务也变得越来越异构和多样化,未来的研究重点包括:

1)开发更有效的图嵌入和图神经网络方法,以便更好地学习动态和异质的图结构与属性信息,从而确保模型的准确性和实用性。

2)设计更高效的图采样方法,以加快图表示学习的训练和推断过程,并尽可能地保持图信息的完整性;同时探索分布式图表示学习方法,以适应大规模分布式图数据的处理需求。

3)继续研究图神经网络的可解释性,以在光网络调度中做出更有利的决策。

参考文献:

- [1] 张宏科,权伟,刘康.算力网络研究与探索[J].中兴通讯技术,2023,29(1):1-5.
- [2] 沈世奎,王光全,唐雄燕.全光算力网络打通算力资源运输的大动脉[J].通信世界,2022(6):31-32.
- [3] BAO B, YANG H, YAO Q, et al. SDFA: a service-driven fragmentation-aware resource allocation in elastic optical networks[J]. IEEE Transactions on Network and Service Management, 2021, 19(1): 353-365.
- [4] ZHAO J, BAO B, YANG H, et al. Holding-time-and impairment-aware shared spectrum allocation in mixed-line-rate elastic optical networks[J]. Journal of Optical Communications and Networking, 2019, 11(6): 322-332.
- [5] BAO B, YANG H, YAO Q, et al. Resource allocation with edge-cloud collaborative traffic prediction in integrated radio and optical networks[J]. IEEE Access, 2023, 11: 7067-7077.
- [6] YU T, YANG H, YAO Q, et al. Multi visual GRU based survivable computing power scheduling in metro optical networks[J]. IEEE Transactions on Network and Service Management, 2024, 21(1): 1302-1315.
- [7] BHATTI U A, TANG H, WU G, et al. Deep learning with graph convolutional networks: An overview and latest applications in computational intelligence[J]. International Journal of Intelligent Systems, 2023(3): 1-28.
- [8] ULLAH I, MANZO M, SHAH M, et al. Graph convolutional networks: analysis, improvements and results[J]. Applied Intelligence, 2022, 52: 1-12.
- [9] HADI M, PAKRAVAN M R, AGRELL E. Dynamic resource allocation in metro elastic optical networks using Lyapunov drift optimization [J]. Journal of Optical Communications and Networking, 2019, 11(6): 250-259.