

基于深度增强学习的软件定义网络路由优化机制

兰巨龙 于倡和* 胡宇翔 李子勇

(国家数字交换系统工程技术研究中心 郑州 450002)

摘要: 为优化软件定义网络(SDN)的路由选路, 该文将深度增强学习原理引入到软件定义网络的路由选路过程, 提出一种基于深度增强学习的路由优化选路机制, 用以削减网络运行时延、提高吞吐量等网络性能, 实现连续时间上的黑盒优化, 减少网络运维成本。此外, 该文通过实验对所提出的路由优化机制进行评估, 实验结果表明, 路由优化机制具有良好的收敛性与有效性, 较传统路由协议可提供更优的路由方案与实现更稳定的性能。

关键词: 软件定义网络; 路由优化; 深度增强学习

中图分类号: TP393

文献标识码: A

文章编号: 1009-5896(2019)11-2669-06

DOI: 10.11999/JEIT180870

A SDN Routing Optimization Mechanism Based on Deep Reinforcement Learning

LAN Julong YU Changhe HU Yuxiang LI Ziyong

(National Digital Switching System Engineering & Technological Research Center, Zhengzhou 450002, China)

Abstract: In order to achieve routing optimization in the Software Defined Network (SDN) environment, deep reinforcement learning is imposed to the SDN routing process and a mechanism based on deep reinforcement learning is proposed to optimize routing. This mechanism can improve network performance such as delay, throughput, and realize black-box optimization in continuous time, which surely reduces network operation and maintenance costs. Besides, the proposed routing optimization mechanism is evaluated through a series of experiments. The experimental results show that the proposed SDN routing optimization mechanism has good convergence and effectiveness, and can provide better routing configurations and performance stability than traditional routing protocols.

Key words: Software Defined Network (SDN); Routing optimization; Deep reinforcement learning

1 引言

软件定义网络(Software Defined Network, SDN)打破了OSI(Open System Interconnection)模型的垂直结构, 实现了控制与转发的分离, 简化了网络的运维管理过程, 为网络的演进与创新提供了更多可能, 被认为是未来网络可能的架构之一。但随着网络控制的日渐精细与网络规模的扩大与应用的不断涌现, 网络流量近乎呈现指数增长, 而用户需求也日趋多样。传统的基于最短路径算法来计算路由的方式, 存在收敛速度慢的问题, 不适合动

态网络, 并且对网络变化的响应也可能导致严重的拥塞。在此背景下, 对SDN网络的路由选路过程进行优化对于保证服务质量, 推动SDN的发展演进, 便显得至关重要。

近些年来, 机器学习是人工智能发展的热门领域, 其在大规模数据处理、分类以及智能决策方面的卓越性能, 使得它可能成为破解目前网路运维与管理僵局的有效武器^[1,2], 特别是为了解决传统选路方式的不足, 许多研究人员尝试将机器学习等诸多智能算法引入SDN路由选路机制, 从而实现路由的智能化、可定制与细粒度管理。

文献[3]基于多种机器学习机制协同的方法提出了一种路由预设计方案, 该方案首先利用恰当的聚类算法如高斯混合模型或者K-mean聚类的方式对数据流进行特征提取, 随后, 利用监督学习机制如极限学习机做流量需求预测。在此基础上, 根据数据流不同约束因子的权重, 提出了一种基于层次分析法的自适应多径路由方法来处理大象流。文献[4-6]

收稿日期: 2018-09-06; 改回日期: 2019-05-12; 网络出版: 2019-05-27

*通信作者: 于倡和 yu_changhe@hotmail.com

基金项目: 国家自然科学基金群体创新项目(61521003), 国家自然科学基金(61502530)

Foundation Items: The National Natural Science Foundation of China for Innovative Research Groups (61521003), The National Natural Science Foundation of China (61502530)

则利用启发式算法,如蚁群算法、遗传算法等来进行选路的优化,但由于启发式算法自身的局限性,所建模型只针对特定问题,并且当网络状态发生变化时,参数可能需要重新调整,可扩展性受限。文献[7,8]都在SDN网络中部署了增强学习模块,其中文献[7]利用QoS感知的奖励函数,通过Q-learning^[9]在SDN网络中实现QoS自适应路由。而文献[8]提出了一种端到端的自适应HTTP流媒体智能传输架构。系统建模基于部分可观察性马尔可夫决策过程,并以Q-learning方法作为最大化QoE的集群决策算法。增强学习在路由选路应用中拥有明显优势,可以实现低时延、高吞吐量和自适应路由,但在传统Q-learning算法中,由于SDN网络对数据流的细粒度管控,维护[状态,动作]→奖励信息的Q表会带来大量的存储资源开销,并且随着Q表规模的扩大,查表的时间开销也是不容忽视的因素。这极大限制了增强学习机制在SDN路由选路中的应用。为了解决这个问题,文献[10]以神经网络替代传统增强学习中的Q表,采用DQN(Deep Q-learning)^[11]的方法来优化选路过程,然而DQN不仅自身不易收敛,且仅能实现离散时间上的控制与优化,与网络的动态特性匹配度不高。

针对上述的现状与挑战,本文将深度增强学习机制DDPG(Deep Deterministic Policy Gradient)^[12]应用于SDN的选路过程,在实现低时延、高吞吐量的同时,有效避免了维护大规模Q表所带来的存储资源与查表的时间开销,实现连续时间上的黑盒优化。并且通过实验测试,验证了DDPG优化SDN路由选路机制不仅自身收敛性与效能良好,且与传统路由算法OSPF(Open Shortest Path First)相比,可以实现更佳的性能指标与稳定性。

2 加装机器学习模块的SDN框架与DDPG一般原理

SDN的转发与控制分离的架构与全局网络视图为网络设计提供了极大的灵活性,并使得基于全局的网络管控与优化成为可能。而机器学习近些年的迅猛发展以及卓越的多维数据处理与智能决策能力,使得其在网络中的应用极富有前景。如何将SDN网络与机器学习算法有机结合,是一个有趣而富有价值的问题^[13]。本节在叙述DDPG技术原理的基础上,提出了一种可行的SDN网络加装机器学习模块的框架,基于此框架,SDN可以实现智慧化的运维。

2.1 加装机器学习机制的SDN网络框架

图1展示集成机器学习的SDN网络的结构,即在SDN控制层面上集成基于机器学习的智能决策

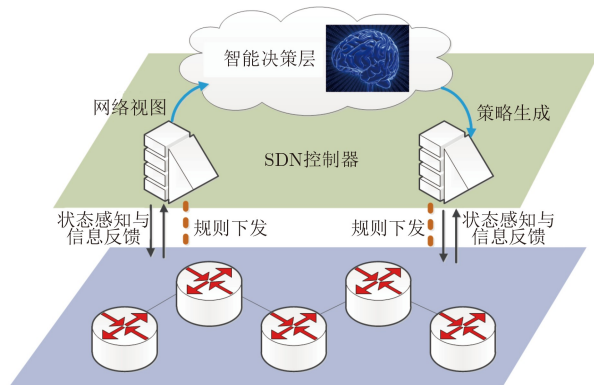


图1 加装机器学习机制的SDN网络架构

层,通过SDN的全网视图与网络测量技术,智能决策层可以实现高效的策略生成,从而做到全局化、实时性、个性化的网络智慧化管控。

具体而言,该框架是在通过SDN架构获取全网视图与网络状态等信息的基础上,通过“智慧决策层”形成相应的网络策略,再借助控制平面生成对应的流表规则进行下发。可以发现,SDN网络的转发与控制相分离的架构以及全网视图能够很好地满足网络智慧化的需求,而机器学习算法的发展为实现轻量而高效的“智慧决策层”提供了契机。网络智慧化的关键在于对网络状态的全局感知以及基于全网信息进行智慧决策,而外界可以通过对智慧决策层的调控来管理网络自身想要的运维性能,而这个调整过程是通过编程实现的。在该架构基础上,可以智慧化地对路由选路、资源适配等一系列运维管理操作进行优化,实现网络空间的“无人驾驶”。

2.2 深度增强学习DDPG的一般原理

增强学习是让智能体在环境中学习动作从而获取最大的奖励值的过程,而深度增强学习是将深度学习的感知能力与增强学习的决策能力结合起来的学习方法。然而一般的基于值的深度增强学习机制如DQN,仅仅针对离散时间的动作控制,无法有效适配连续动作的建模与控制,针对动态性与实时性较强的网络系统,其控制策略会存在误差与延时。而可以实现连续时间控制与优化的基于策略的增强学习方法DPG(Deterministic Policy Gradient)^[14],则受困于线性函数拟合的策略函数精确性较差与过拟合。针对以上问题,DeepMind团队提出的深度增强学习模型DDPG将DQN方法与DPG方法以actor-critic^[12]框架结合起来,并以神经网络来拟合策略函数与Q函数,形成一种更加高效且稳定的离散动作控制模型,整体框架流程如图2所示。其中, μ 与 Q 分别代表神经网络拟合的确定性策略函数与Q函数。

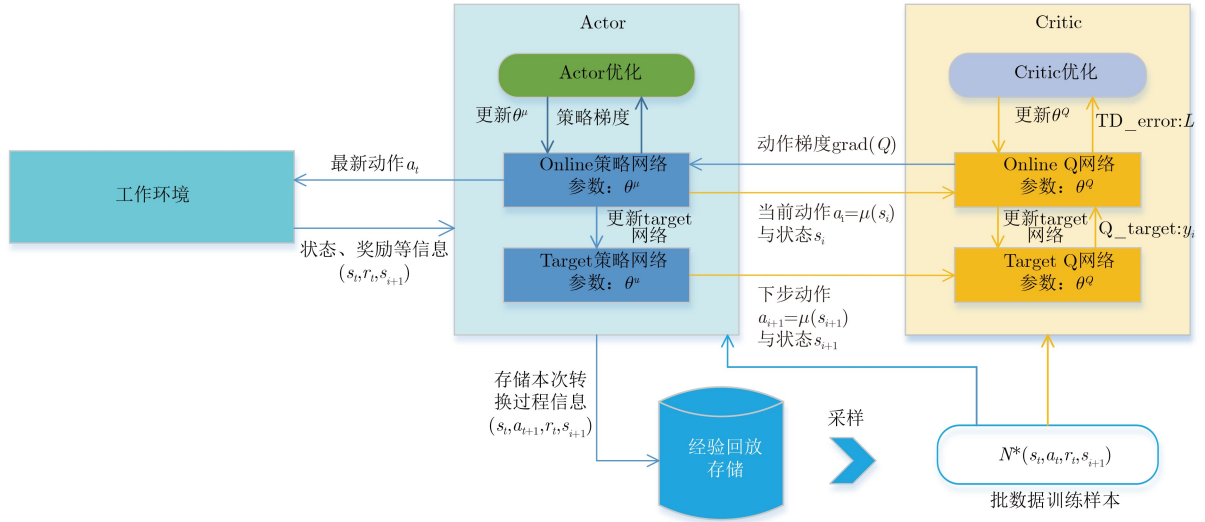


图2 DDPG的训练运行框架

在DDPG的actor-critic架构中，Actor模块采用DPG方法而Critic模块采用DQN方法，而每个模块都由两个神经网络组成，分别为用于训练学习的Online网络与用于阻断训练数据相关性的Target网络，即Target网络是相同结构但参数不同的Online网络，其参数为一段时间前的Online网络参数，即Online网络会定期将自身参数传递给Target网络用于更新。在一轮训练回合中，每次与环境的互动的转换过程信息都被存储到经验池中，并且，神经网络的学习样本是从经验池中采样抽取若干个转换过程信息构成的，各模块的Online网络通过学习训练样本来不停地更新神经网络参数。

DDPG的参数更新过程有Actor模块用于拟合DPG策略函数的神经网络更新与Critic模块的DQN神经网络更新组成，但二者的更新相互渗透结合，彼此相关。对于Actor模块的DPG神经网络，其更新需求反向传递的是策略梯度，策略梯度更新如式(1)所示

$$\begin{aligned} \nabla_{\theta^\mu} J &= \text{grad}[Q] * \text{grad}[\mu] \\ &\approx \frac{1}{N} \sum_i \nabla_a Q(s, a | \theta^Q) |_{s=s_i, a=\mu(s_i)} \nabla_{\theta^\mu} \mu(s | \theta^\mu) |_{s_i} \end{aligned} \quad (1)$$

其中， J 为衡量策略 μ 表现的性能函数。前半部分 $\text{grad}[Q]$ 为来自Critic网络的动作梯度，用于表征Actor的动作获得更大回报的移动方向，而后半部分 $\text{grad}[\mu]$ 为来自Actor网络的参数梯度，用于表征Actor的神经网络应该如何调整自身参数，才能使得神经网络以更高的可能性选取这个高回报动作。两者相乘结合，意义为Actor模块的神经网络朝着有更高可能获取高回报的方向修改自身参数。

对于Critic模块DQN网络的更新，其更新反向传递是TD_error，其更新过程如式(2)所示

$$L = \frac{1}{N} \sum_i (y_i - Q(s_i, a_i | \theta^Q))^2 \quad (2)$$

即Critic模块的TD_error等于Online网络的Q值与Target网络的Q值均方误差。其中， y_i 为Target网络的基于下一步状态 s_{i+1} 与动作 a_{i+1} 的Q值，值得注意的是，这里的下一步动作来自Actor模块的Target网络，即该Target网络用段时间前Online网络传递来的参数 μ' 来与下一步状态 s_{i+1} 选择出下一步的动作 a_{i+1} ，而 y_i 的计算如式(3)所示

$$y_i = r_i + \gamma Q'(s_{i+1}, \mu'(s_{i+1} | \theta^{\mu'}) | \theta^Q) \quad (3)$$

对于Online网络的Q值，则为对当前状态下Actor网络选取的当前动作的评估，即Actor模块的Online网络将自身选取的当前动作 a_i 与传递给Critic模块的Online网络用于评估并生成Q值。

3 DDPG机制优化SDN路由框架与方法

在本节中，本文所设计的DDPG优化SDN路由选路的框架与机制运行流程进行介绍，通过运行该机制，SDN网络可以实现对定制性能参数的优化，如最小化时延、最小化跳数、最大化吞吐量等。这与所选取的网络控制策略与网络运维的目的有关。利用DDPG优化SDN选路，可以实现对网络的连续时间的实时控制，并且是一种黑盒优化，实现自动地优化选路，有效缓解运维压力。

在2.2节中已对DDPG的基本原理与部署机器学习的SDN框架进行了介绍。本节将DDPG部署到SDN中，利用其实现对SDN网络路由的优化，基本框架流程如图3所示。DDPG智能体通过3个信号

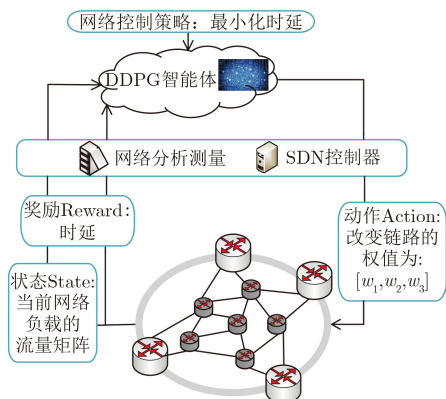


图3 DDPG优化SDN路由选路的框架设计

即状态、动作、奖励与环境进行交互。其中，状态 s 为当前网络负载的流量矩阵(Traffic Matrix, TM)，而智能体对环境所采取的动作 a 为改变网络链路的权值，即通过改变网络链路的权值来改变所涉源目的节点对的数据流所采用的路径。最后，智能体的奖励 r 与所采取的网络运维策略有关，可以是时延、吞吐量等单个性能参数，当然也可以是均衡多个参数的综合策略。这些控制策略的调整可通过更改Reward的设置来实现。而后文将以最小化时延作为控制策略举例说明。

DDPG智能体的训练目的为根据输入状态 s 来找到最优动作 a 以最大化回报 r ，而DDPG优化SDN路由选路的一般过程，可简述为：通过网络分析测量与SDN控制器，DDPG智能体可以获取准确的网络状态 s ，而后DDPG智能体可以确定一个优化动作，即一套链路权值 $[w_1, w_2, \dots, w_n]$ 。根据该套权值配置重新计算数据流链路，并通过集中控制的SDN控制器实现规则下发。随后，通过网络测量分析获取新路由实施后的奖励 r 与新的网络状态，如此循环实现对所需网络性能的优化。

4 实验评估

本节对DDPG优化路由机制的性能进行评估测试，并主要关注DDPG智能体自身的有效性与收敛性和DDPG智能体与其余路由方法相比较下的性能优势。测试实验运行的硬件环境NVIDIA Tesla P100 GPU与24GB的DDR4 内存，操作系统为Ubuntu 16.04，所采用的机器学习框架为Keras+TensorFlow，并用OMNeT++搭建仿真网络环境^[15]。在同一构造拓扑下，本文对DDPG智能体自身收敛性与有效性进行测试，并将DDPG优化路由性能与当前主流的路由协议OSPF与大量随机生成路由配置进行比较，4.1节中详细介绍了实验方法，而4.2节就实验结果进行分析比较。

4.1 实验方案

设计实验采取OMNeT++构造网络环境，实验拓扑采用Sprint网络的主干网络拓扑，由25个节点和53个链路组成，并假设链路带宽相同。为逼近实际网络场景，实验中设置了多个不同流量强度，分别占全网总带宽的不同百分比。对于每个流量强度，在相同流量总量下，分别使用重力模型^[16]生成若干个不同的流量矩阵，即生成了多个差异化流量总量与分布的流量矩阵，用于训练与测试等。

为了验证DDPG优化路由机制自身的有效性与收敛性，将DDPG智能体在不同流量强度等级下训练不同的步数，通过对训练后智能体进行性能测试来验证其训练有效性与收敛性。而为了验证DDPG优化SDN路由选路的性能优势，共设计了2套实验方案。分别为：(1)在不同流量强度下，比较DDPG的优化路由与大量随机生成的路由配置的性能，将DDPG优化路由的性能与随机生成的100000个有效路由配置进行比较，这些随机路由配置在本场景有效在于皆可实现无环路与所有节点相互可达，且通过大数量来保证对比测试数据集的代表性与有效性。(2)在相同网络环境下，假设拓扑各链路的权值相同，比较OSPF协议与DDPG的选路性能。

4.2 实验结果

在本文实验中，以SDN运维策略为最小化时延为例，将其作为DDPG优化路由的性能测度，在DDPG优化SDN路由机制有效的前提下，其余控制策略可通过定制reward来类似实现，故此处不予一一列举。而在实验中，关注2点，即DDPG自身的收敛性与有效性以及DDPG优化SDN路由机制与OSPF协议、随机路由配置在最小化时延方面的性能对比。

4.2.1 DDPG优化SDN路由机制的收敛性与有效性

在实验中，采用4种不同等级的流量强度，分别占全网带宽的10%、40%、70%与100%，对于每个流量强度，分别生成250个流量矩阵。随后，DDPG智能体在不同流量强度下用所对应的生成流量矩阵分别训练2000、4000、10000、20000、50000、80000、100000次后，用1000个流量矩阵作为输入，测试DDPG智能体性能，得到训练后智能体所选取的路由方案。值得注意的是，给定流量矩阵与路由方案后，OMNeT++仿真环境可获取网络时延等网络性能参数作为DDPG路由优化机制的奖励。为了保证数据有效性，用测试流量矩阵进行的1000次测试的测试结果取平均值，作为反应训练效果的测度。实验结果如图4所示。

实验结果表明，DDPG智能体通过训练，可以

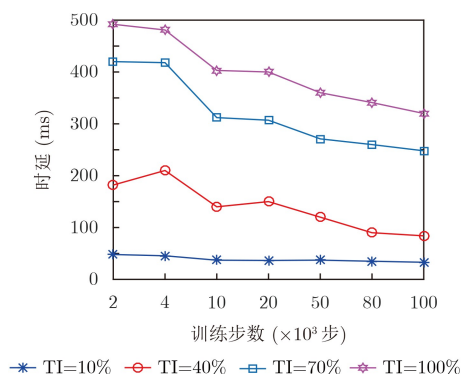


图4 不同流量强度下网络的时延随训练步数的变化

有效削减网络的时延，并且随着训练步数的增加，网络的时延也不断削减优化，证明DDPG智能体优化网络的性能不断提升。如图4中，在70%网络容量的流量强度下，DDPG智能体训练100000步的优化性能比只训练2000步的优化性能提升了40.4%，这说明所提出的DDPG优化SDN路由选路方案具有明显的实效，且收敛性良好。

4.2.2 DDPG优化SDN路由机制的性能

为了验证DDPG优化SDN路由选路的性能优势，实验分别将DDPG智能体与大量随机生成路由配置与OSPF路由算法进行性能对比。

首先将随机生成的100000个路由配置实施到各流量强度下的测试流量矩阵，从而通过OMNeT++得到相应的时延。这里将训练100000步后的DDPG智能体作为测试比较的对象，对比随机路由的时延与DDPG路由优化机制生成的路由策略的时延。对比结果以箱线图的形式在图5中表示。

图5中矩形模块的上部与底部分别代表实验所得时延观测数据的上四分位数与下四分位数值，并且，矩形中间横线代表实验观测数据的中位数。而由矩形延伸出来的直线的上端与下端分别表示观测数据在内限范围中的最大值与最小值，此外，出于简洁，内限范围外的无效数据不予显示。实验结果表明，绝大多数DDPG优化出的路由配置的时延都

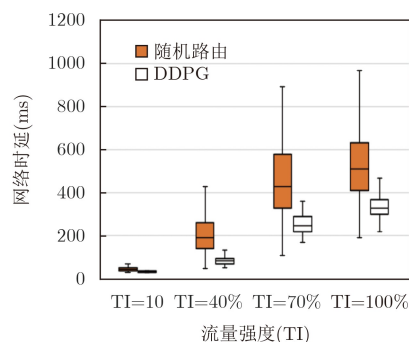


图5 DDPG智能体与随机路由对比

小于大量随机生成路由时延的下四分位值，且DDPG智能体的优化路由配置所得的最小时延都十分逼近大量随机路由由测试集所得的最小时延结果。以上结果充分验证了DDPG优化SDN路由机制的性能优越性与有效性。

随后，在流量强度相同且实验拓扑的链路权值设为同一值的条件下，分别运行DDPG智能体与OSPF来作为路由选路机制，此处运行的DDPG智能体经过了100000次的训练。将DDPG优化SDN路由机制与传统的OSPF协议的运行情况进行对比，对比结果如图6所示。

图6中展示的为两种路由选路方式运行过程中，8000个数据包传输的时延情况，其中，图6(a)为OSPF协议运行时，数据包的时延情况，可以发现，大部分OSPF协议传输的数据包，其时延峰值超过300 ms。而DDPG优化路由的运行时延中，如图6(b)，只有少部分时延峰值超过300 ms，说明经过DDPG优化的路由选路方案可以实现比传统OSPF协议更加优秀的性能，证明了本文DDPG优化SDN路由选路机制的有效性与性能优势。

最后，关于DDPG智能体的训练与决策时延，深度增强学习机制是优化网络运维管理过程的有效工具，但在其训练收敛之前，智能体的训练开销较大与优化性能不佳，特别是当网络中存在大量网络服务与应用时。每当网络中出现新的服务时，智能

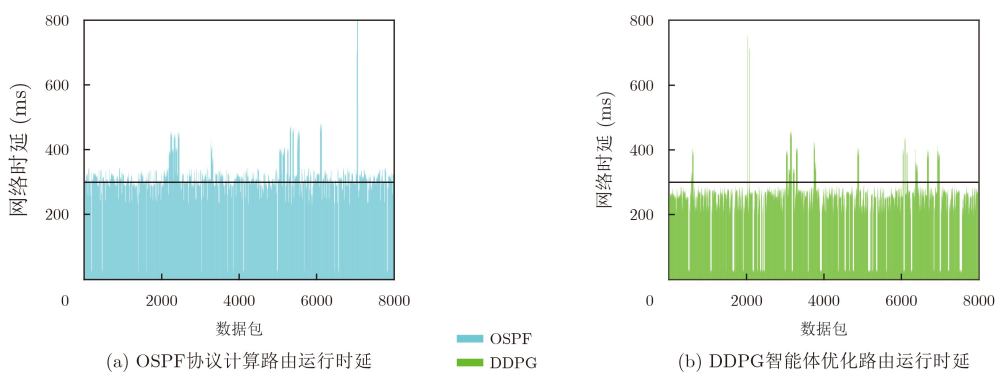


图6 DDPG与OSPF的网络运行时延对比

体都需要重新训练收敛,这限制了智能体优化路由选路的灵活性与可扩展性。如何寻找有效机制来削减智能体的训练开销与周期,仍是一个开放性问题。目前可能的研究方向或解决办法有:(1)通过离线学习收敛后再线上运行;(2)增加先验知识以加快训练过程,如转移学习、模仿学习^[17]。

但值得注意的是,一旦智能体训练收敛后,便不需要随着所输入的网络状态的变化而再次训练收敛,并可以通过矩阵乘法一步计算出近似最优解,其决策时间复杂度约为 $O(n^2)$,其中 n 代表网络中边数。相比之下,启发式算法需要采取许多步骤来收敛到一个新的结果,这导致较高的计算时间成本。例如,蚁群算法的时间复杂度为 $O(n(n-1)mt)$,其中, n 代表网络中边数, m 为蚂蚁数量, t 为迭代次数。因此,DDPG智能体可以有效削减决策时延,这对于动态网络的实时控制有较大的意义。

5 结束语

本文提出了一种可行的SDN加装机器学习机制的运维架构,并将DDPG深度增强学习机制加装到SDN网络中,用于优化SDN网络的路由选路过程,实现连续时间的黑盒优化,简化网络运维过程。此外,本文对DDPG优化SDN路由选路机制的性能进行了实验测试。实验结果表明,本文提出的SDN路由优化机制具有良好的收敛性与有效性,与传统路由协议相比较,能提供更加稳定且优越的选路功能,对于SDN网络的运维管理,具有较大的实际意义与价值。

参 考 文 献

- [1] BOUTABA R, SALAHUDDIN M A, LIMAM N, *et al.* A comprehensive survey on machine learning for networking: Evolution, applications and research opportunities[J]. *Journal of Internet Services and Applications*, 2018, 9(1): 16. doi: [10.1186/s13174-018-0087-2](#).
- [2] FADLULLAH Z M, TANG Fengxiao, MAO Bomin, *et al.* State-of-the-art deep learning: Evolving machine intelligence toward tomorrow's intelligent network traffic control systems[J]. *IEEE Communications Surveys & Tutorials*, 2017, 19(4): 2432–2455. doi: [10.1109/COMST.2017.2707140](#).
- [3] LI Wei, LI Guojun, and YU Xiufen. A fast traffic classification method based on SDN network[C]. The 4th International Conference on Electronics, Communications and Networks, Beijing, China, 2015: 223–229.
- [4] WANG Fu, LIU Bo, ZHANG Lijia, *et al.* Dynamic routing and spectrum assignment based on multilayer virtual topology and ant colony optimization in elastic software-defined optical networks[J]. *Optical Engineering*, 2017, 56(7): 076111. doi: [10.1117/1.OE.56.7.076111](#).
- [5] PARSAEI M R, MOHAMMADI R, and JAVIDAN R. A new adaptive traffic engineering method for telesurgery using ACO algorithm over Software Defined Networks[J]. *European Research in Telemedicine*, 2017, 6(3/4): 173–180. doi: [10.1016/j.eurtel.2017.10.003](#).
- [6] WANG Junchao, DE LAAT C, and ZHAO Zhiming. QoS-aware virtual SDN network planning[C]. 2017 IFIP/IEEE Symposium on Integrated Network and Service Management, Lisbon, Portugal, 2017: 644–647. doi: [10.23919/INM.2017.7987350](#).
- [7] LIN S C, AKYILDIZ I F, WANG Pu, *et al.* QoS-aware adaptive routing in multi-layer hierarchical software defined networks: a reinforcement learning approach[C]. 2016 IEEE International Conference on Services Computing, San Francisco, USA, 2016: 25–33. doi: [10.1109/SCC.2016.12](#).
- [8] JIANG Jingyan, HU Liang, HAO Pingting, *et al.* Q-FDBA: Improving QoE fairness for video streaming[J]. *Multimedia Tools and Applications*, 2018, 77(9): 10787–10806. doi: [10.1007/s11042-017-4917-1](#).
- [9] SUTTON R S and BARTO A G. Reinforcement Learning: An Introduction[M]. Cambridge, MA: The MIT Press, 1988.
- [10] SENDRA S, REGO A, LLORET J, *et al.* Including artificial intelligence in a routing protocol using Software Defined Networks[C]. 2017 IEEE International Conference on Communications Workshops, Paris, France, 2017: 670–674. doi: [10.1109/ICCW.2017.7962735](#).
- [11] MNIH V, KAVUKCUOGLU K, SILVER D, *et al.* Human-level control through deep reinforcement learning[J]. *Nature*, 2015, 518(7540): 529–533. doi: [10.1038/nature14236](#).
- [12] LILLICRAP T P, HUNT J J, PRITZEL A, *et al.* Continuous control with deep reinforcement learning[P]. USA, Patent, WO2017019555, 2017.
- [13] MESTRES A, RODRIGUEZ-NATAL A, CARNER J, *et al.* Knowledge-defined networking[J]. *ACM SIGCOMM Computer Communication Review*, 2017, 47(3): 2–10. doi: [10.1145/3138808.3138810](#).
- [14] SILVER D, LEVER G, HEESS N, *et al.* Deterministic policy gradient algorithms[C]. International Conference on Machine Learning, Beijing, China, 2014: I-387–I-395.
- [15] VARGA A and HORNIG R. An overview of the OMNeT++ simulation environment[C]. The 1st International Conference on Simulation Tools and Techniques for Communications, Networks and Systems & Workshops, Marseille, France, 2008: 60.
- [16] ROUGHAN M. Simplifying the synthesis of internet traffic matrices[J]. *ACM SIGCOMM Computer Communication Review*, 2005, 35(5): 93–96. doi: [10.1145/1096536.1096551](#).
- [17] PAN S J and YANG Qiang. A survey on transfer learning[J]. *IEEE Transactions on Knowledge and Data Engineering*, 2010, 22(10): 1345–1359. doi: [10.1109/TKDE.2009.191](#).

兰巨龙: 男, 1962年生, 教授, 博士生导师, 主要研究方向为新型网络体系结构与网络安全。

于倡和: 男, 1993年生, 硕士, 研究方向为新型网络体系结构与网络安全。