

# IT and Multi-layer Online Resource Allocation and Offline Planning in Metropolitan Networks

Miquel Garrich, José-Luis Romero-Gázquez, Francisco-Javier Moreno-Muro, Manuel Hernández-Bastida, María-Victoria Bueno Delgado, Anderson Bravalheri, Navdeep Uniyal, Abubakar Siddique Muqaddas, Reza Nejabati, Ramon Casellas, Óscar González de Dios, Pablo Pavón Mariño

**Abstract**—Metropolitan networks are undergoing a major technological breakthrough leveraging the capabilities of software-defined networking (SDN) and network function virtualization (NFV). NFV permits the deployment of virtualized network functions (VNFs) on commodity hardware appliances which can be combined with SDN flexibility and programmability of the network infrastructure. SDN/NFV-enabled networks require decision-making in two time scales: short-term *online resource allocation* and mid-to-long term *offline planning*. In this paper, we first tackle the dimensioning of SDN/NFV-enabled metropolitan networks paying special attention to the role that latency plays in the capacity planning. We focus on a specific use-case: the metropolitan network that covers the Murcia - Alicante Spanish regions. Then, we propose a latency-aware multilayer service-chain allocation (LA-ML-SCA) algorithm to explore a range of maximum latency requirements and their impact on the resources for dimensioning the metropolitan network. We observe that design costs increase for low latency requirements as more data center facilities need to be spread to get closer to the network edge, reducing the economies of scale on the IT infrastructure. Subsequently, we review our recent joint computation of multi-site VNF placement and multilayer resource allocation in the deployment of a network service in a metro network. Specifically, a set of subroutines contained in LA-ML-SCA are experimentally validated in a network optimization-as-a-service architecture that assists an Open-Source MANO instance, virtual infrastructure managers and WAN controllers in a metro network test-bed.

**Index Terms**— Metropolitan Networks; Network Function Virtualization, Optimization; Software-Defined Networking

## I. INTRODUCTION

THE infrastructure in telecom operators' central offices (COs) has notably evolved in recent years with advancements in software-defined networking (SDN) [2] and network function virtualization (NFV) [3]. SDN and NFV are enabling unprecedented capabilities jointly employing programmable networking and the deployment of virtualized

network functions (VNFs) on commodity hardware appliances. In fact, a network service chain, defined as a path coupled to traverse a certain sequence of VNFs given by a network service, unifies the concepts of SDN and NFV. These technologies shook the telecommunications landscape with profound implications both in (i) the way networks are operated, i.e. short-term actions, and (ii) mid-to-long term network dimensioning and planning decisions. In particular, numerous initiatives leverage on SDN and NFV concepts for pursuing an efficient use of the IT and network infrastructure [4]. Encompassing the technological breakthrough of SDN/NFV in metropolitan networks, the European project Metro-Haul [5] proposes scalable, dynamic and efficient metro networks to efficiently interface 5G access and high-capacity core/backbone networks. The nucleus of the Metro Haul proposal is the dynamic interconnectivity of two different types of nodes, access metro edge node (AMEN) and metro core edge node (MCEN), both with computational capabilities to permit SDN/NFV functionalities to provision VNFs close to the end users.

In this context, SDN and NFV are considered two major technology enablers for realizing 5G networks [6]. In fact, it is essential to maintain high-quality service performance to end users, which are relevant to permit novel applications emerging from Industry 4.0, vehicular networks, tactile Internet and Internet of Things (IoT), 6-DoF (degree of freedom) virtual reality. Special attention must be given to latency requirements in these services, being critical for their performance, successful deployment and usage. Both SDN/NFV capabilities and latency requirements are relevant for the advent of the 5G-era networks because they drive requirements at the control/management and operational levels [7]. Indeed, SDN/NFV are essential to enable “a *scalable management framework enabling fast deployment of novel applications*”

This work was supported in part by the Spanish Government: ONOFRE-2 project under Grant TEC2017-84423-C3-2-P (MINECO/AEI/FEDER, UE) and the Go2Edge project under Grant RED2018-102585-T; and by the European Commission: METRO-HAUL project (G.A. 761727) and INSPIRING-SNI project (G.A. 750611). Preliminary results contained in this paper were presented at ECOC 2019 [1].

M. Garrich, J.-L. Romero-Gázquez, F.-J. Moreno-Muro, M. Hernández-Bastida, M.-V. Bueno Delgado and P. Pavón Mariño are with Telecommunication Networks Engineering Group, Department of Information and Communication Technologies, Universidad Politécnica de Cartagena; M.-V. Bueno Delgado and P. Pavón Mariño are also with E-lighthouse Network Solutions, Cartagena 30202, Spain (e-mail: miquel.garrich@upct.es,

josel.romero@upct.es, javier.moreno@upct.es, manuel.hernandez@upct.es, mvictoria.bueno@upct.es; pablo.pavon@upct.es).

A. Bravalheri, A. S. Muqaddas, N. Uniyal, and R. Nejabati, are with High Performance Networks Group, Department of Electrical and Electronic Engineering, University of Bristol, Bristol, BS81UB United Kingdom (e-mail: a.bravalheri@bristol.ac.uk, abubakar.muqaddas@bristol.ac.uk, navdeep.uniyal@bristol.ac.uk, reza.nejabati@bristol.ac.uk).

R. Casellas is with the Centre Tecnològic de Telecomunicacions de Catalunya (CTTC/CERCA), Castelldefels 08860, Spain (e-mail: ramon.casellas@cttc.es).

Ó. González de Dios is with Telefónica GCTO, Madrid, 28050 Spain (e-mail: oscar.gonzalezdedios@telefonica.com).

while end-to-end latency needs to be ensured below 1 ms for selected flows, among other 5G key performance indicators (KPIs) to be accomplished [7].

In this paper, we consider both online short-term resource provisioning actions and mid-to-long term offline planning/dimensioning decisions of SDN/NFV-enabled metropolitan networks paying special attention to the latency role in the latter case. We begin with a network dimensioning use case that covers the Murcia - Alicante Spanish regions. We use our recently proposed NFV over IP over WDM (NIW) library, an open-source modelling and evaluation framework for SDN/NFV metropolitan networks proposed within the Metro-Haul context. In particular, NIW is a library added to Net2Plan [8] open-source network planning software, specifically targeting the provision, design and evaluation of SDN/NFV networks. Leveraging NIW, we propose a latency-aware multilayer service chain allocation (LA-ML-SCA) algorithm, used to explore a range of latency requisites, and we observe its impact on the IT resources dimensioning. Results clearly show that design costs increase as the proportion of low-latency flows grow, as more NFV capabilities need to get closer to the network edge, and thus the economies of scale are degraded. Then, we review our recent NFV-IP-WDM resource provisioning demonstration [1], which comprises a full-interconnected SDN-NFV test-bed with externalized intelligence hosted in a Net2plan-based network *Optimization as a Service (OaaS)* architecture [9]. Specifically, several subroutines contained in LA-ML-SCA to iteratively provision network resources (i.e. with the goal of network planning) are used to assist ETSI Open-Source MANO (OSM) [10] in multi-virtual infrastructure managers (VIMs), multi-VNFs network service instantiation and multi-layer WAN resource allocation with a unified approach.

The remainder of the paper is organized as follows. Section II reviews recent related works. Section III describes modelling and design assumptions of the network planning use case. In Section IV, we perform the offline network dimensioning that evaluates the latency vs. cost trade-off considering services' maximum permitted latency to reach the first VNF of the service chain. Section V reviews our recent online resource allocation demonstration. Section VI concludes the paper.

## II. RELATED WORK

SDN and NFV are two technologies that have attracted a lot of interest from the industry and the academia in recent years. Here, we review related work on SDN and NFV that comprises from the optical layer up to the IT resources in the nodes' location from two different perspectives: in-operation resource allocation for network control/orchestration and offline planning for network dimensioning that consider latency requirements.

### A. In-operation (online) SDN/NFV Resource Allocation

Recent works provide the SDN/NFV technological solutions for in-operation resource allocation describing network control/orchestration architectures and protocols. For instance, the work in [11] proposes an architectural framework for orchestrating mobile access networks over a multi-layer

network that exploits the benefits of SDN and NFV. Subsequently, Casellas *et al.* focus on the control aspects in control, management, and orchestration systems for optical networks that consider multi-layer optical networks and cloud resources for enabling 5G network slicing [12]. Recently, works in [13] and [14] focus on wide-area networks (WANs) for applying SDN and NFV, hence enabling the IT-network joint perspective in multi-layer networks. Both [13] and [14] are based on well-known protocols (e.g. PCEP) to propose an ad-hoc architecture that targets latency aspects in [14]. Remarkably, Bravalheri *et al.* [15] report an extension of OSM that specifically targets the WAN resource allocation jointly with multi-site IT infrastructure management.

An architectural alternative is considered in the proofs of concept demonstrated in [16] and [17]. In particular, an optimization engine based on Net2Plan enters in the scene so as to efficiently employ IT and network infrastructure in a joint manner leveraging on algorithms commonly used for planning purposes. Specifically, [16] targets latency requirements in the SDN/NFV infrastructure and [17] further elaborates on the IT resources in the operators' COs.

### B. Offline SDN/NFV Network Planning

#### 1) Latency Modelling and Assumptions.

Among the reviewed literature, several works assume latency models for VNFs following a queue approach. The work in [18] is based on a G/G/m queue model for intra-DC service chains. In this context, [19] considers a latency model at the datacenters (DCs) that accounts on an intra-datacenter CPU resource sharing model among the instantiated VNFs in the DC and NIC queues. Besides, authors in [19] disseminate radio, fixed and aggregation-core network link congestion in end-to-end latency calculations. Few papers consider the impact of the processing overhead as the load increases in each VNF. The VNF execution time in [20] depends on an approximation of the Amdahl's Law considering the used processors executing the VNF. Also, authors are aware of the end-to-end latency. Furthermore, Savi *et al.* [21], assume end-to-end latency for services while considering CPU load balancing to calculate VNF traversing time. Besides, authors provide the VNFs sequence of the service chain to satisfy the proposed services.

A different approach is considered in [22], where the end-to-end delay is obtained with two methods, (i) the *standard* method, where latency is a function of the aggregation of the VNF bit-rate, and is sensitive to the traffic load, and (ii) the *fastpath* method, that considers a constant forwarding latency. In [23], authors consider the latency as the end-to-end delay, with the VNF processing time expressed as function of the working frequency of the processor (CPU).

#### 2) IP over WDM requirements in NFV chaining

The related works reviewed above motivate the need to focus on service chain allocation (and VNF placement) considering the IP over WDM performance to meet the requirements of the most stringent 5G use cases. In this context, authors in [24] propose an algorithm which performs dynamic service-chaining in an optical metro-area network which jointly minimizes the average number of nodes required to host VNF instances as well as the blocking probability. Although optical-

electrical-optical (OEO) conversion and forward error correction (FEC) overhead penalties are mentioned in [24], no specific values are provided nor their impact is deeply analyzed in the joint NFV placing and IP over WDM resource allocation. Relevantly, Pedro *et al.* [25] analyze the performance of three metropolitan node architectures in terms of the number of transceivers, capacity and latency for a 5, 10 and 15-node metropolitan ring topology. Latency results in the range of 1 to 3 ms, which consider propagation and OEO conversion, highlight the importance of network design in these scenarios [25]. Indeed, FEC penalties range from 50  $\mu$ s [26] up to 150  $\mu$ s [27], which combined with OEO conversion are commonly assumed to jointly cause 100  $\mu$ s of delay [25], [28].

### C. Our Contribution

In this work, the novelty relies on the usage of the same framework based on Net2plan for a twofold contribution: (i) an offline network dimensioning with a cost model and a realistic topology description, and (ii) an online resource allocation based on an experimental validation. Note that other/pervious work address these aspects separately. Regarding offline network dimensioning, section II.B highlights the need for an analysis on the cost implications when performing VNF placement jointly with IP over WDM resource allocation considering propagation time and FEC plus OEO conversion penalty. We report a realistic use-case of offline network dimensioning in sections III and IV. Relevant for online resource allocation, our work reported in section V represents a novel concept in which Net2Plan assists both OSM and WAN manager with the joint decision of VNF placement and network resource allocation.

## III. NFV-IP-WDM USE CASE: MURCIA-ALICANTE REGION

In this section, we first describe the topological characteristics of the metropolitan network produced for this use case. Then, we detail the traffic, latency, cost and service models used.

### A. Topology Description

Fig. 1 illustrates the network topology under study, which is located in the South-East of Spain. 71 nodes (31 nodes in Murcia and 40 in Alicante) correspond to cities and towns considering realistic information from a nation-wide telecom operator list of metropolitan COs. The nodes are connected via a realistic WDM topology, synthesized in a form that matches the topological characteristics of reference metro networks provided in Metro-Haul project. In particular, we consider three reference networks from Telecom Italia (TIM) operator [29] with three sizes (see Table I).

#### 1) Node categories

We define three node categories: AMEN, MCEN without connectivity to the core network and MCENb in which 'b' denotes connectivity to the backbone network. The category of each node in the Murcia-Alicante (M-A) topology is defined following a similar ratio of categories in the three instances of the TIM topology (see Table I). The result is that the 71 nodes in M-A are distributed into 62 AMEN, 7 MCEN and 2 MCENb.

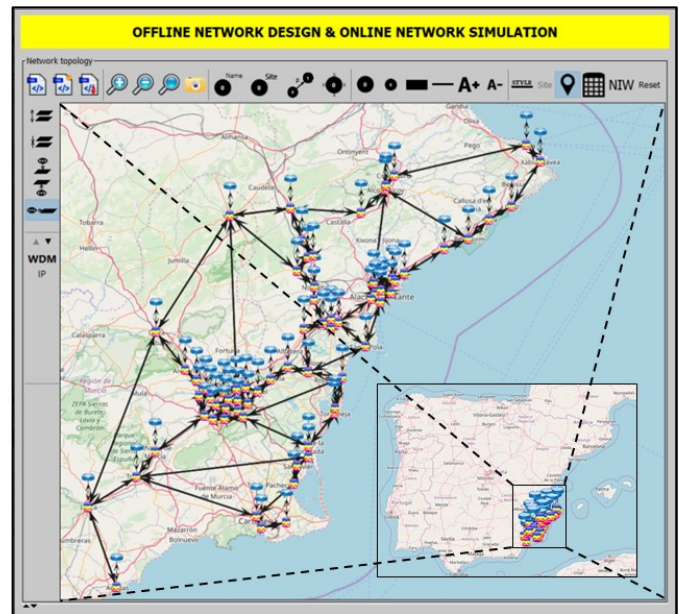


Fig. 1. Topology of Murcia-Alicante Spanish regions under study.

Then, the MCENb are assigned to Murcia1 and Alicante1 nodes which are at the center of the regions' capital cities. Finally, an additional MCEN is assigned to the capitals (Murcia and Alicante) and the remainder 5 MCEN are distributed taking into account the most populated municipalities while configuring a homogeneous node type distribution across the network.

#### 2) Number of fibers and origin/destination

One hundred and two (102) bidirectional fiber links are set to reach an average node degree of 2.87 (i.e.  $102 \times 2 / 71$ ). Subsequently, we jointly consider the combination of node degree distribution (see Fig. 2) and network diameter (see Table I) to define the origin/destination node of each fiber.

The node degree distribution for M-A is similar to the three instances of the TIM networks. The WDM plant determines the network diameter expressed as the longest shortest path in number of hops, between all possible source-destination pairs. Although small and medium TIM network instances are similar to M-A in terms of average node degree and fibers, the network diameter diverges due to the particular geolocation and distribution of nodes. Indeed, TIM networks present a homogeneous distribution of the nodal degree whereas M-A concentrates the high degree nodes in two central city areas and spreads low degree nodes toward the network edge. Thus, the trade-off between average node degree and network diameter has been assessed to form a realistic M-A network.

TABLE I  
TOPOLOGY CHARACTERISTICS OF TIM AND M-A NETWORKS

| Tab                 | TIM small  | TIM med    | TIM large   | M-A (ours) |
|---------------------|------------|------------|-------------|------------|
| Links               | 72         | 143        | 219         | 102        |
| Nodes               | 52         | 102        | 129         | 71         |
| AMEN                | 44 (84.6%) | 89 (87.2%) | 113 (87.6%) | 62 (87.3%) |
| MCEN                | 6 (11.5%)  | 11 (10.8%) | 11 (8.5%)   | 7 (9.9%)   |
| MCENb               | 2 (3.9%)   | 2 (2%)     | 5 (3.9%)    | 2 (2.8%)   |
| Average node degree | 2.77       | 2.80       | 2.75        | 2.87       |
| Net. diameter*      | 9          | 12         | 15          | 13         |

AMEN: Access Core Edge Node, MCEN(c): Merto-Core Edge Node (connected to the core network).

\* Network diameter: Max number of hops between all possible node pairs.

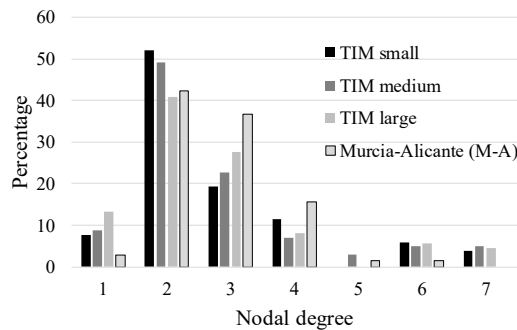


Fig. 2. Nodal degree distribution of three TIM and M-A networks.

### B. Traffic Model

Here, we first describe how the population of the municipalities is assigned to each network node. Then, we detail our traffic model.

#### 1) Population assigned to each node

The assignment of population to each network node in the NIW framework is performed according to the three following cases. First, the population of a municipality with a single node representation is directly assigned to the `node.population` characteristic that defines that node. Second, in case of a municipality with two or more node representations in the framework, its population is uniformly distributed among those nodes. Third, the population of a municipality with no assigned node in the framework is added to the nearest node. Table II lists the number of municipalities and inhabitants in Murcia [30] and Alicante [31] regions and our assignment of population into adjacent nodes without node are included (ordered from highest to lowest population) into the nearest nodes until 95 % of population is assigned.

#### 2) Traffic generation model

Users in each node are supposed to be using 50 different services, each one reflected by two service chain requests, upstream and downstream. This sums 3550 bidirectional service chain requests (71 nodes, by 50 services). Each service chain has a number of VNF instances randomly chosen between 1 and 5. This modelling decision is based on the work reported in [24], which defines the characteristics for 6 realistic service chains (with their number of VNFs): cloud gaming (5), augmented reality (5), VoIP (5), video streaming (5), massive IoT (3) and smart factory (2). All VNFs require 2 CPU cores, 4 GB of RAM and 20 GB of HDD. No traffic compression is considered in any VNF. We further assume that each VNF instance (and its associated resources) is devoted exclusively to the service chain it belongs.

Upstream service chain starts in a particular node and ends in the nearest MCENb. The traffic intensity is proportional to the node population, and to a random uniformly picked multiplicative factor in the range [0.1, 1.0]. The traffic is normalized so that the maximum traffic per service injected among all services and all nodes in the network, equals 1 Gbps.

### C. Latency Assumptions and Requirements

Two major latency contributions are considered: propagation delay in the optical fibers and OEO conversion penalty. The propagation delay ( $d_f$ ) in the optical fiber is calculated as  $d_f = \frac{c}{n} \cdot l_f$ , where  $c$  is the speed of light ( $3 \cdot 10^5$  km/s),  $n = 1.5$

TABLE II

POPULATION CHARACTERISTICS OF THE MURCIA AND ALICANTE REGIONS

|                                           | Alicante  | Murcia    |
|-------------------------------------------|-----------|-----------|
| Municipalities                            | 141       | 45        |
| Municipalities with ( $\geq 1$ ) node     | 30        | 19        |
| Municipalities to cover 95% of population | 60        | 30        |
| Population (in 2017)*                     | 1,825,332 | 1,463,249 |
| Population covered (95%)                  | 1,734,066 | 1,390,087 |
| Population not considered                 | 91,266    | 73,162    |

\* Governmental data of Murcia [30] and Alicante [31] regions.

is the fiber's refractive index and  $l_f$  is the fiber link length in km. We consider a constant OEO conversion time of 0.2 ms at each OEO hop, which includes the FEC penalty. This value is chosen as upper bound (i.e. worst-case) that guarantees the OEO/FEC overhead penalties (see section II.B).

In this work, latency requirements are limited to a maximum delay between the user and the first VNF in the service chain. In each test, such latency requirements are random uniformly picked per service chain request between  $[0.5, L_{\max}]$  milliseconds. Latency limits are enforced for reaching just the first VNF because the delay experienced to traverse the first (as any) VNF largely depends on software processing performance as reviewed in section II. However, the LA-ML-SCA algorithm will instantiate all the VNFs (not just the first) in the same AMEN or MCENb nodes for a given service. This is a worst-case scenario from the point of view of cost: a latency limit requiring instantiating the first VNF close to the user, results in all the VNFs in the chain being instantiated in the same data center, while in a less restrictive scenario those could be left for a more centralized data center (e.g. in the MCENb). Additionally, this modelling assumption is made to perform VNF placement decisions that aim at satisfying end-to-end low-latency requirements as the ones described in [24]. In particular, [24] describes service chains for VoIP, augmented reality and smart factory with latency requirements of 5 ms, 1 ms and 1 ms, respectively. Consequently, such low values of latency can only be ensured by placing all VNFs close to the user.

### D. Cost Model

In this work, we concentrate on the costs of the IT resources required for VNF instantiation. In particular, we are interested in capturing how the distribution of IT resources in micro-data centers, instead e.g. of its concentration in one or few data centers per metro network, can impact the overall cost. To evaluate this, we consider a cost model that reflects the effects of the economy of scale. Such economic laws have been shown to apply to all aspects of economy, and mean that the cost-per-unit-produced decreases as the number of units produced increases. In our case, means i.e. that building a data center of size  $2Z$  in a node, costs less than building two data centers of size  $Z$  in two nodes.

There are a number of mathematical relations that have been shown to model the economies of scale effects. All of them are characterized by a *concave* (sublinear) curve describing the cost of a system, respect to its size (see [32] for details). In this paper, we choose the square-root law, which appears in many economic scenarios as an accurate descriptor of the economies of scale law. According to this, we estimate the cost of a data center placed in a particular node as:

---

**Algorithm 1.** Latency-aware multilayer service chain allocation

---

**Input:**  $G = (N, E)$ ,  $L_{\max}$

**Output:** VNF placement, IP links and routing, lightpaths

**Being:**  $G$  graph,  $N$  nodes,  $E$  fibers,  $e$  fiber,  $L_{\max}$  latency limit,  $l(p(r))$  latency of path  $p$  of service chain request  $r$ ,  $V(r)$  VNFs of  $r$ .

**Begin:**

1. Sort service chain requests in decreasing traffic
  2. Create empty list of established lighpaths:  $LP$
  3. **foreach** ( $r \in SCR$ ) **do**:
  4. Calculate shortest path  $p(r)$  in km from  $oriNode(r)$  to the nearest core node ( $destNode(r)$ )
  5. **while** ( $l_{dest}(p(r)) > l(r)$ )
  6. Calculate  $l_{dest}(p(r))$  // latency to  $destNode(r)$  with OEO conversion penalty
  7. **if** ( $l_{dest}(p(r)) \leq l(r)$ )
  8. Establish lightpath:  $\lambda \forall e \in p(r) \mid \lambda(e) \notin LP$
  9. Instantiate  $V(r)$  in  $destNode(r)$
  10. **else**
  11. **foreach**  $\lambda_{inter} \in LP$ (ending at  $destNode(r)$ ) **do**:
  12. Calculate  $l_{2-hop}(p(r)) = l(\lambda_{ori}^*) + l(\lambda_{inter})$
  13. **if** ( $l_{2-hop}(p(r)) \leq l(r)$ )
  14. Establish lightpath to  $interNode(r)$ :  $\lambda_{ori}$
  15. Instantiate  $V(r)$  in  $destNode(r)$
  16. **break**
  17. **end if**
  18. **end foreach**
  19. **if** ( $l(\lambda_{direct}^*(r)) \leq l(r)$ )
  20. Establish lightpath to  $destNode(r)$ :  $\lambda_{direct}$
  21. Instantiate  $V(r)$  in  $destNode(r)$
  22. **break**
  23. **end if**
  24. Calculate  $l_{pNode}(r)$  // Calculate latency to the previous node  $pNode(r)$  with OEO conversion in  $p(r) = p(r) - 1$
  25. Update  $destNode(r) = pNode(r)$
  26. **end if**
  27. **end while**
  28. **end foreach**
- 

$$C_{node} = C \sqrt{c_1 \cdot CPU + c_2 \cdot RAM + c_3 \cdot HDD}.$$

$CPU$ ,  $RAM$  and  $HDD$  are the total amount of CPUs (assumed all of the same type), RAM (in GB) and HDD (in GB) required in the data center, respectively. Coefficients  $c_1$ ,  $c_2$  and  $c_3$  weight the contribution of each resource to the total deployment of the NFV-enabled CO. In particular, we consider a common server size composed of 2 CPU cores, 4GB of RAM and 2,000 GB of HDD, in which the cost of these components is 70 €, 50 € and 85 €, respectively. This permits to set  $c_1 = 70/2$  (CPU coefficient),  $c_2 = 50/4$  (RAM coefficient) and  $c_3 = 85/2000$  (HDD coefficient). Factor  $C$  is set to 1, reflecting that we are interested in comparing normalized cost values among the different alternatives.

#### E. The NFV over IP over WDM (NIW) Library

In this use case, we use the NFV over IP over WDM (NIW) open-source library [33], with available documentation in Javadoc format [34]. NIW is currently contained in Net2Plan and shipped within its latest release [35]. NIW is developed in the context of the Metro-Haul project [5] with the purpose of modeling SDN/NFV-enabled networks in a versatile framework which, based on its Excel spreadsheet importer and reporting functionalities, permits the abstraction of the underlying representations for rapid network engineering (network capacity planning and dimensioning) [33].

#### IV. NIW MURCIA-ALICANTE USE-CASE EVALUATION

In this section, we perform a dimensioning study of the metropolitan networks that cover Murcia and Alicante regions. We present and describe the proposed LA-ML-SCA algorithm. Then, the dimensioning study executes LA-ML-SCA and computes the resulting CO and network costs. Finally, simulation results are presented and discussed.

##### A. LA-ML-SCA Algorithm

The LA-ML-SCA algorithm is shown in Algorithm 1. receives as an input (i) an IP over WDM network modeled as a graph  $G$  with  $N$  nodes and  $E$  fibers, (ii) with available IT resources in the COs, (iii) the service chain requests together with their latency requisites. Its goal is satisfying the requests by allocating service chains, creating the needed VNF instances and IP links that are realized via lightpaths over the WDM plant. Latency-aware means that allocations are restricted to satisfy the latency requisites, as well as of course avoid IT resource oversubscription and spectrum clashing.

The LA-ML-SCA algorithm first sorts the service chain request  $r$  according to its injected traffic (higher traffic first) in line 1. Then, it sequentially processes each request  $r$  one by one (line 3) by calculating its shortest path  $p(r)$  from the origin node to the nearest core node (line 4). The core of the LA-ML-SCA algorithm (lines 5-27) is a **while** loop that aims at allocating all the VNFs of  $r$  in the furthest node from the user (closest to the core, and thus more centralized). This potential allocation is evaluated in three approaches. Firstly, lines 7-9 evaluate if latency requirements are satisfied with  $p(r)$  considering OEO and propagation. If so, the algorithm establishes a lightpath per link in  $p(r)$ , which creates a new logical IP link or increases IP capacity (in case an IP link is present). Secondly, lines 11-18 aim to reuse previously established lightpaths ( $\lambda_{inter}$ ) evaluating potential lighpaths ( $\lambda_{ori}^*$ ) in terms of total latency with two OEO conversion. Thirdly, lines 19-23 evaluate if a direct lighpath satisfies the latency requirements.

Finally, the LA-ML-SCA algorithm iterates the three approaches above reducing one hop the shortest path in each iteration (lines 24-25). In case any approach above is satisfied, all VNFs are instantiated in the current destination node of  $p(r)$ . We recall that all VNFs of  $r$  are instantiated in the same node for ensuring that stringent latency constraints are satisfied. Additionally, this assumption is made to avoid the trivial solution in which only the first VNF of  $r$  is placed in a node close to the user, while all the other VNFs are consolidated in few nodes, without caring about end-to-end latency. The output of the LA-ML-SCA algorithm is used to compute the resources employed at all COs, number of lightpaths (i.e. transponders) and IP switching occurrences.

The LA-ML-SCA algorithm has a polynomial complexity. Its computational time is dominated by the shortest-path calculations which are performed for each service chain request  $r$  once in line 4 and *iteratively* three times in lines 8, 14 and 20. Note that the number of iterations that the latter three shortest-path calculations are performed depends on the latency requirements. Specifically, in case the latency requirements are not satisfied, the core **while** loop in the LA-ML-SCA

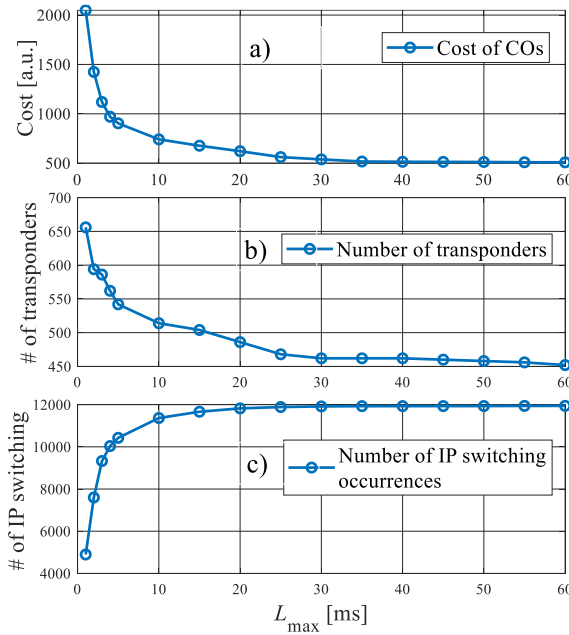


Fig. 3. a) CO's cost, b) number of transponders and c) number of IP switching occurrences as function of the maximum latency allowed to reach the first VNF.

algorithm iterates these three calculations with a one-hop reduction from the destination node on each iteration. Consequently, high computational time is expected for low latency requirements. These iterations are in the order of  $N$  (worst case), whereas the shortest-path calculations are implemented with a complexity  $O(E+N \log N)$  and exhibit a more uniform execution time.

### B. Simulation Results and Discussion

We consider a M-A network infrastructure with enough pool of IT resources (CPU, RAM, HP) in all COs to place all service chain requests without contention. The network size is obtained executing the proposed LA-ML-SCA algorithm, which places all VNFs over the M-A network. Note that LA-ML-SCA progressively establishes IP links (and lightpaths) to meet the services' latency requirements, starting with an empty network.

We investigate the network design cost (in arbitrary units see section III.D) as a function of services' latency requirements. In particular, we explore a range of  $L_{\max}$  latency values (maximum latency from the user to the first VNF) in a range between 1 ms up to 60 ms, which correspond to the most stringent and the most permissive scenarios, respectively. We recall that  $L_{\max}$  is used to define the maximum latency constraints for each service chain request, which are uniformly picked between  $[0.5, L_{\max}]$ .

Fig. 3a) shows the cost of COs as a function of the considered values of  $L_{\max}$ . Results show an abrupt cost increase for values below 10 ms. The cost of COs to guarantee the services' latency of 1 ms reaches up to four times the value of more permissive scenarios  $L_{\max} \geq 30$  ms. Indeed, low latency values are only attainable instantiating VNFs close to the origin nodes of the service chains, which decentralizes the IT resources into a larger amount of COs. Conversely, more permissive scenarios (i.e.  $L_{\max} \geq 30$  ms) tend to concentrate the IT resources into MCENb, which leveraging on the economy of scale, reduces the COs' cost. Fig. 3b) reports the number of transponders used

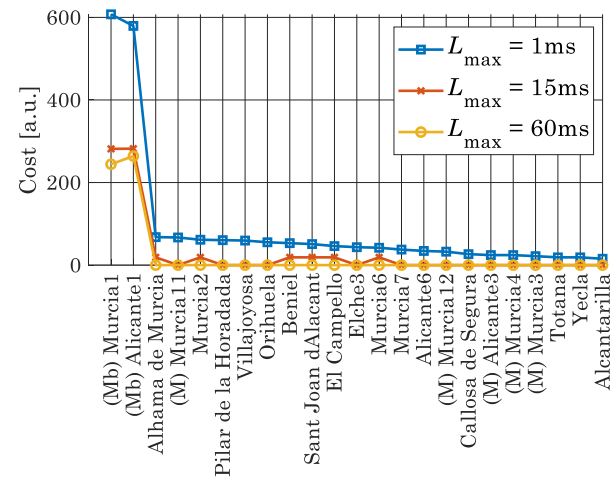


Fig. 4. CO costs for maximum latencies allowed of 1, 15 and 60 ms.

in each network dimensioning also as a function of the considered values of  $L_{\max}$ , which exhibits a similar trend as the cost of the COs. This behavior occurs because LA-ML-SCA tends to deploy direct lightpaths from the end user to the nearest MCENb (last node of the service chain) in case the latency requirements are low, hence optical bypass is required at intermediate nodes. Fig. 3c) shows the total number of IP switching occurrences performed by all the 3550 service chain requests. Remarkably, IP switching has the opposite trend compared to both cost of the COs and number of transponders, because IP switching is largely employed for high-latency cases. In fact, high latency requirements correspond to permissive cases and thus OEO conversions are allowed through the service chain. In summary, a trade-off is evidenced because low-latency requirements devote capital expenditures on IT resources at COs toward the edge jointly with transponders, whereas high-latency requirements migrate cost towards IP equipment.

Fig. 4 illustrates the per node CO cost distribution considering  $L_{\max}$  values of 1 ms, 15 ms and 60 ms. Nodes in the x-axis are sorted in decreasing order according to their cost for  $L_{\max} = 1$  ms, denoting MCEN and MCENb as (M) and (Mb), respectively. The MCENb central nodes in the capital cities (Murcia1 and Alicante1) concentrate the largest amount of resources. For  $L_{\max} = 1$  ms, Murcia1 and Alicante1 present a 10-fold CO cost value compared to the remainder nodes. Additionally, the resources not allocated in MCENb nodes are distributed among various AMEN and few MCEN.  $L_{\max}$  values of 15 ms and 60 ms further concentrate the CO cost in the two MCENb nodes, up to no resources needed in both MCEN and

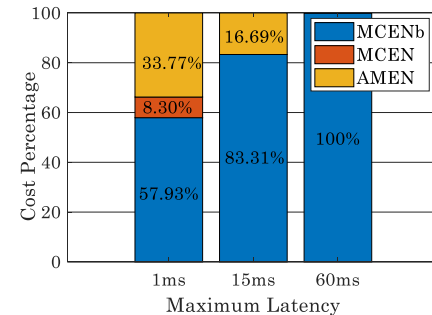


Fig. 5. CO cost distribution between MCENb, MCEN and AMEN.

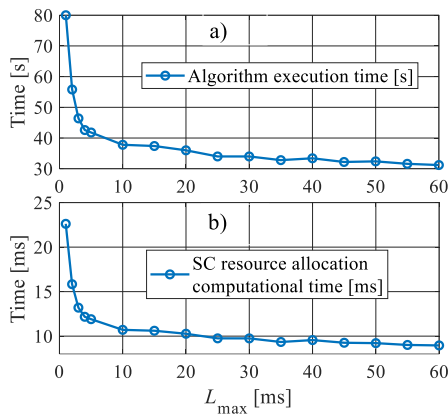


Fig. 6 a) Total LA-ML-SCA algorithm computational time and b) SC request resource allocation computational time as function of the maximum latency allowed to reach the first VNF.

AMEN nodes in the latter case. Relevant to this observation, Fig. 5 groups the CO cost distribution per node type for the three  $L_{\max}$  values under study. Indeed, as anticipated in Figs. 4 and 5, CO cost is concentrated in MCENb nodes for high (i.e. more permissive)  $L_{\max}$  values whereas low-latency values spread the CO cost towards nodes not connected to the backbone network.

Fig. 6a) and b) report the computational times of the entire LA-ML-SCA algorithm and for the allocation decision of each service chain request, respectively, both as a function of the maximum latency allowed to reach the first VNF. As anticipated at the end of section IV.A, high computational times are experienced for low-latency requirements. In particular, the core `while` loop of the LA-ML-SCA is iteratively executed for satisfying stringent cases of low latency. Finally, it is worthwhile mentioning that the total execution time in Fig. 6a) corresponds to  $3550 \times$  (i.e. number of service chain requests) the computational time employed for the allocation decision of each service chain request reported in Fig. 6b).

## V. NFV-IP-WDM EXPERIMENTAL DEMONSTRATION

In this section, we review our recently proposed architecture, workflow and demonstration in which Net2Plan performs a network Optimization-as-a-Service (OaaS) in the joint computation of multi-site VNF placement and multilayer resource allocation in the deployment of a network service in a metro network [1]. Specifically, the online resource allocation performed herein includes several specific subroutines within LA-ML-SCA to assist an Open-Source MANO instance and WAN Infrastructure Manager(s) (WIM) in an experimental proof of concept comprising a metro-network test-bed.

### A. SDN/NFV Architecture for Wide Area Networks

Fig. 7 shows the components comprised in the architecture for the experimental validation of IT over IP over WDM resource allocation. Firstly, on top of the functional schema, the role of the operator is represented by an Operations Support System (OSS) realized with a specific extension of Net2plan. The cloud/network optimization tasks are externalized in the network OaaS module [9], which relies on optimization procedures on top of network controllers and orchestrators

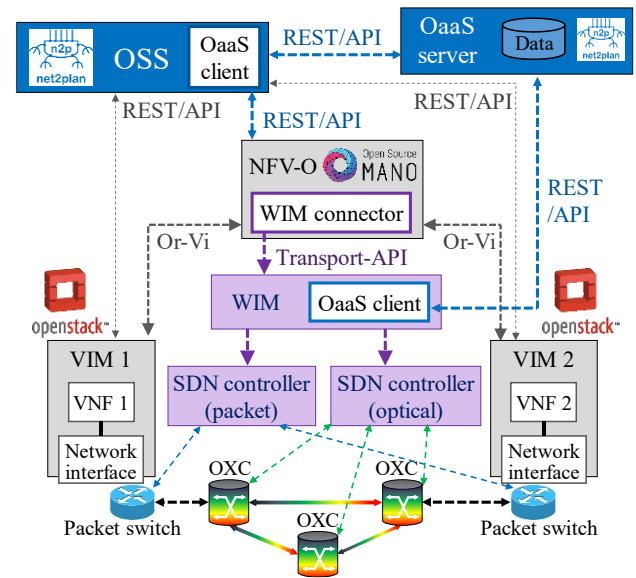


Fig. 7. Functional Architecture

exploiting a software layer of third-party applications to request, in this case, efficiency in IT and network resource allocations. The Net2Plan-based OaaS module exposes a REST-API, which includes calls to request the computation of a joint VNF placement and multilayer resource allocation solution for a network service request. The VNF instantiation is communicated to OSM, while the multilayer allocation is stored in a database, and is later retrieved by a API call from the WIM.

Secondly, the control plane is composed of:

- i) ETSI's NFV-Orchestrator Open-Source MANO (OSM), in charge of orchestrating network services over IP over optical metro networks and the VNFs instantiation.
- ii) Two VIMs, implemented with OpenStack clusters, which are responsible for the virtualization infrastructure by hosting the virtual machines (VMs) of the VNFs.
- iii) The SDN-control is split in an SDN controller per layer, one devoted to the packet layer and the other one in charge of the optical layer.
- iv) The WIM hierarchically on top of the SDN controllers. Especially relevant for SDN-NFV WAN-based work is the presence of the WIM [15], because it links OSM with the SDN-controllers, via Transport-API. Additionally, the WIM interconnects remote VNFs, belonging to the same network service, instantiated in different VIMs along the multi-layer transport network in this WAN environment. Moreover, the WIM includes an OaaS client to request optimal paths from the OaaS server.

Finally, the data plane considers a multi-layer transport network, where the packet layer is represented by packet switches over the optical layer that interconnect the compute nodes under the control of the VIM. The devices, both packet and optical layer, are abstracted by the T-API based WIM for configuration purposes.

### B. Workflow of the WAN-based Network Service Instantiation

To efficiently use both IT and network resources when performing the network service deployment is of importance in the context of the Metro-Haul project. To this end, the proposed

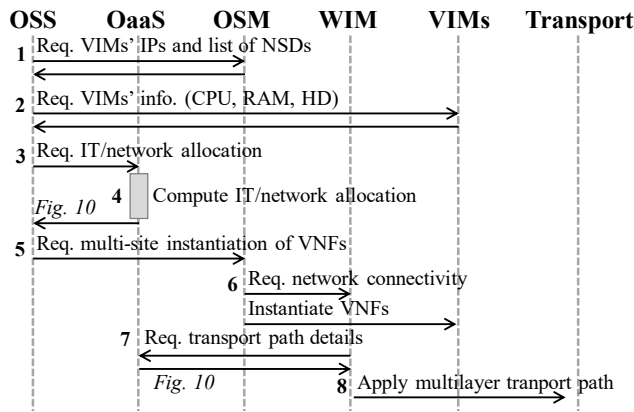


Fig. 8. Simplified schema of the workflow.

architecture combines the Metro-Haul control, orchestration and management (COM) service platform with an external optimization engine to allow efficient use of IT/network resources. The proposed architecture operates with the following workflow (also illustrated in Fig. 8):

1. The OSS, which already contains the multi-layer transport topology in an .n2p (Net2Plan) file, queries OSM for the VIM IP addresses and the Network Services (NSs) available to be instantiated via OSM REST API requests.
2. Once the IP addresses of the VIMs are known by the OSS, it retrieves all the IT information (CPU, RAM, HDD) of the compute nodes and updates the .n2p file.
3. The OSS defines service chain requests regarding the network service to be deployed and it is included in the .n2p

file, ready to be sent via REST/API to the OaaS server as inputs of the algorithm.

4. The OaaS server receives the inputs and the algorithm jointly solves (i) the VNFs placement (in the VIMs), (ii) the allocation of the flows between VNFs over the multilayer network. Relevant for (i), NIW's method `wNet.addVnfInstance` corresponds to `netPlan.addResource`, which is used in this experimental validation. Similarly, for (ii) the network flow computation uses NIW library's functionalities `wNet.getKShortestServiceChainInIpLayer` and `wNet.getKShortestWdmPath`, which are comprised within Net2Plan's generic method `GraphUtils.getKMinimumCostServiceChains`. More details regarding NIW's functionalities can be found in [33]. Additionally, in terms of computational time, the variation of the LA-ML-SCA algorithm here employed is executed in 4.96 ms, which corresponds to half of the lowest times reported in Fig. 6b). This is because the topology employed in this experimental implementation is limited to 9 nodes and 13 fibers. After the OaaS computation, the VNF placement is returned to the OSS, whilst the calculated network allocations are stored in an internal database.
5. Then, the OSS instance commands OSM to instantiate the VNFs in the VIMs selected by the algorithm hosted in the OaaS server.
6. OSM informs the WIM about the multi-VIM connectivity request.

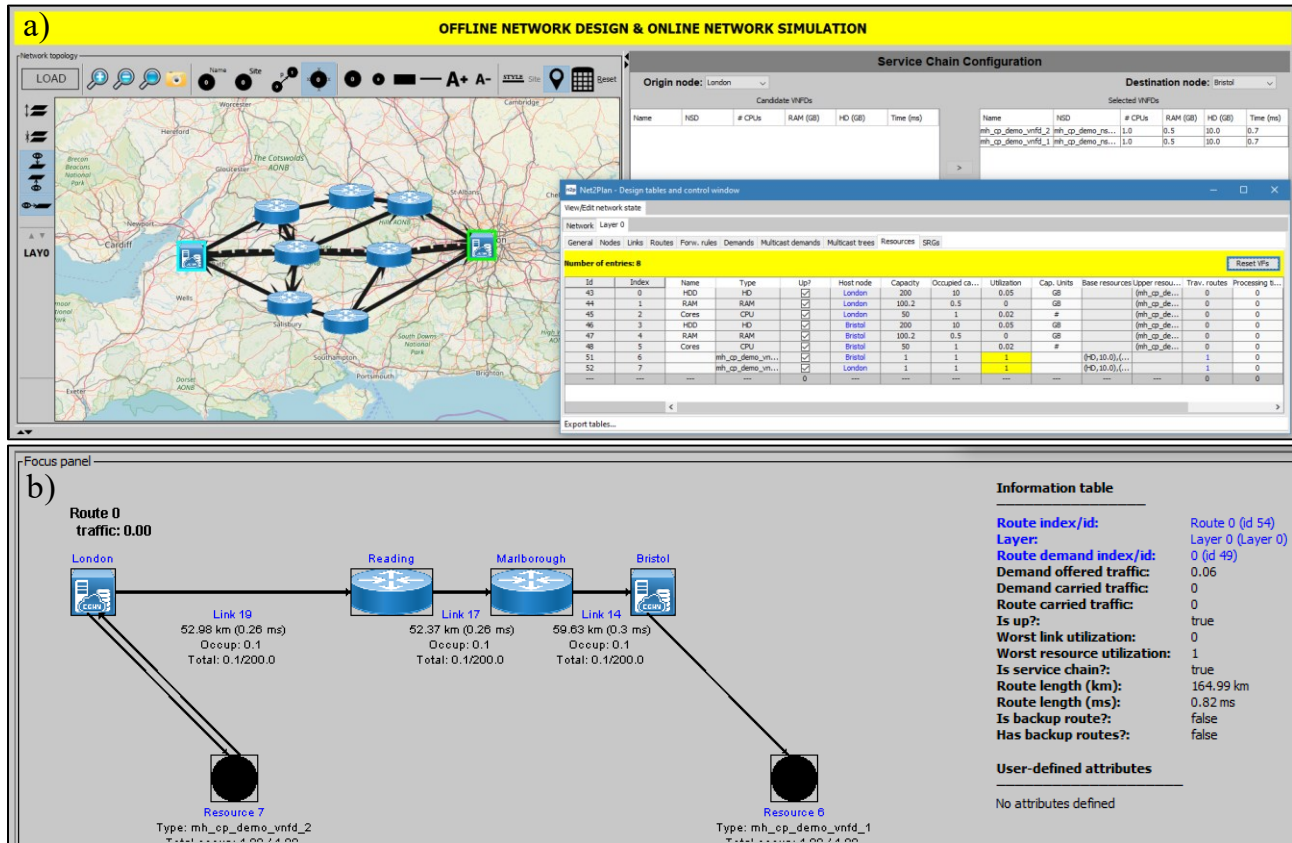


Fig. 9. Net2Plan (OSS) GUI with the network service. a) Topology view and characteristics of the VNFs, and b) network service detail.

7. The WIM queries the OaaS module for the previously calculated path in the transport stored in the database.
8. The WIM contacts both packet and optical SDN controllers to configure the transport devices in order to provide connectivity between the VIMs enabling the VNFs connection; thus satisfying the network service deployment.

### C. Testbed Configuration and Experimental Outcomes

The experimental setup to implement this SDN-NFV WAN architecture proposal is performed in two testbed islands, the two Net2plan instances (OSS and OaaS) are located in a laboratory in Cartagena (Spain) whilst the VIMs, the ETSI OSM, WIM, SDN controllers and the multi-layer topology are placed at the High Performance Networks, University of Bristol (UK). Both testbeds are connected by a private VPN to provide control-plane visibility. The multi-layer infrastructure consists of two OpenStack datacenters interconnected by a combination of packet switches and optical cross connects (OXC) which are SDN-enabled and managed by a combination of SDN controllers as in the functional schema shown in Fig. 7.

Fig. 9 shows Net2plan's GUI with the topology that assumes one datacenter placed in Bristol and the other one in London; where the optical network between the datacenters emulates a realistic WAN domain. Fig. 9b) displays the service chain solution, corresponding to the instantiation of a two-VNF network service solved by the OaaS module (in step 4 in Fig. 8) with the VNF placement and the path in the multi-layer

transport network. Note that Fig. 9b) depicts the information to be stored in the database after the optimization task, origin and destination nodes and the set of links that compose the path traversing the VNFs in the correct order. Note that latency values are computed also in this experimental validation in a per-link basis (0.26 ms, 0.26 ms and 0.3 ms) and for the entire route length (0.82 ms) as shown in Fig. 9b). These values could be used to validate also in this instance of Net2Plan the subroutines employed in the offline network dimensioning part of the paper (section IV). Nonetheless, the experimental validation here reported targets the verification of the COM service platform for successfully performing the service chain instantiation comprising VNFs in multiple VIMs and network path.

Fig. 10 shows a capture of the REST/API GET response regarding used both in the queries to the OaaS performed from the OSS and the WIM. The JSON file includes the network service name and VNF placement instantiation details for the OSS (again, used in step 4 in Fig. 8) and link IDs of the path for the multilayer resource allocation to be carried out by the WIM. After processing the reply from the OaaS module, the WIM oversees the inter-VNF connectivity over the multi-layer infrastructure.

## VI. CONCLUSION

This paper covered two aspects of SDN/NFV-enabled metropolitan networks: short-term resource allocation actions and mid-to-long term planning and dimensioning decisions. We began with the latter aspect with the capacity planning of a metropolitan network considering the impact in the IT infrastructural cost that latency requirements may impose, by forcing the placement of decentralized data centers closer to the user. To this end, we employed Net2Plan's NIW library in a specific metropolitan network (Murcia and Alicante Spanish regions) proposing a Latency-Aware Multilayer Service Chain Allocation (LA-ML-SCA) algorithm. We analyzed the impact of latency limitations on IT resources deployment. The cost model employed captures the economies of scale laws, reflecting the extra costs coming from decentralizing data centers. The obtained results showed that design costs increase for low latency requirements as NFV capabilities need to get closer to the network edge. Indeed, a drastic increase in the overall cost was observed for latency limits below 10 ms. Relevant for the SDN/NFV resource allocation aspect, we reviewed our recent experimental validation in which Net2Plan plays a network OaaS role to support OSM, VIMs and WAN infrastructure manager on the decision of the NFV placement and multilayer resources to be employed.

## REFERENCES

- [1] F. J. Moreno-Muro, M. Garrich, C. San-Nicolás-Martínez, M. Hernández-Bastida, P. Pavón-Mariño, A. Bravalheri, A. Muqaddas, N. Uniyal, R. Nejabati, R. Casellas, and O. González de Dios, "Joint VNF and multilayer resource allocation with an open-source optimization-as-a-service integration," *European Conference on Optical Communication (ECOC)*, Dublin, Ireland, Sept. 2019.
- [2] D. Kreutz, F. M. V. Ramos, P. E. Verissimo, C. E. Rothenberg, S. Azodolmolky and S. Uhlig, "Software-Defined Networking: A Comprehensive Survey," in *Proceedings of the IEEE*, vol. 103, no. 1, pp. 14-76, Jan. 2015.

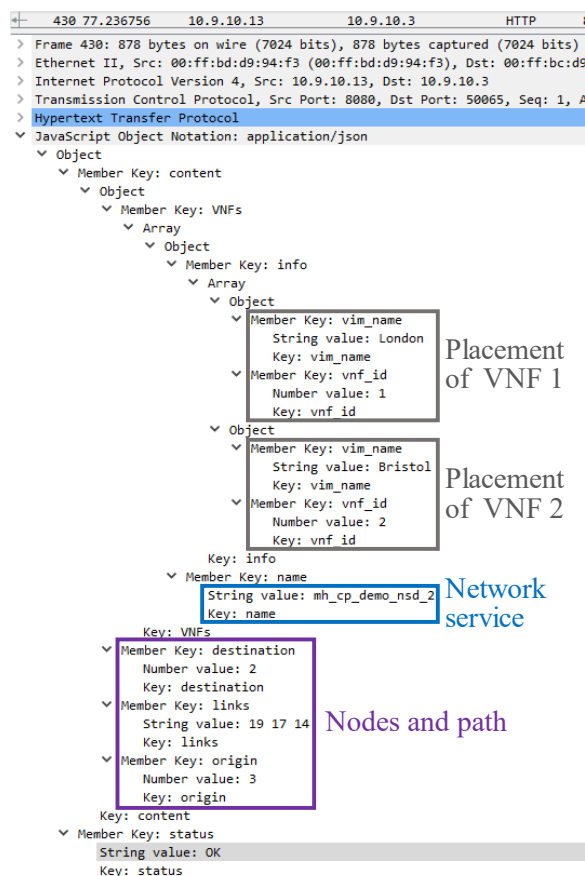


Fig. 10. REST-based response to a GET query capture used in two interactions: OaaS with OSS for IT resources (step 4, Fig. 8); and OaaS with WIM for network resources (step 7, Fig. 8).

- [3] M. Chiosi, *et al.* "Network functions virtualisation: An introduction, benefits, enablers, challenges and call for action," in *Proc. SDN OpenFlow World Congress*, Darmstadt, Germany, Oct. 2012, pp. 22–24. [Online]. Available: [http://portal.etsi.org/NFV/NFV\\_White\\_Paper.pdf](http://portal.etsi.org/NFV/NFV_White_Paper.pdf). Accessed on: Nov. 11, 2019.
- [4] M. Garrich, F. J. Moreno-Muro, M. V. Bueno-Delgado, and P. Pavón-Mariño, "Open-source Network Optimization Software in the Open SDN/NFV Transport Ecosystem," *Journal of Lightwave Technology*, vol. 37, no. 1, Jan. 2019.
- [5] Metro-Haul project. [Online]. Available: <https://metro-haul.eu/>. Accessed on: Nov. 11, 2019.
- [6] F. Z. Yousaf, M. Bredel, S. Schaller, and F. Schneider, "NFV and SDN - Key Technology Enablers for 5G Networks," *Journal on Selected Areas in Communications*, vol. 35, no. 11, pp. 2468–2478, Nov. 2017.
- [7] 5G infrastructure public private partnership (5G PPP) key performance indicators (KPIs). [Online]. Available: <https://5g-ppp.eu/kpis/>. Accessed on: Nov. 11, 2019.
- [8] Net2Plan. [Online]. Available: <http://www.net2plan.com/>. Accessed on: Nov. 11, 2019.
- [9] M. Garrich, C. San-Nicolás, F. J. Moreno-Muro, M. V. Bueno-Delgado, P. Pavón-Mariño, "Network Optimization as a Service with Net2plan", *European Conference on Networks and Communications (EuCNC)*, Valencia, Spain, June 2019.
- [10] Open Source NFV Management and Orchestration (MANO). [Online]. Available: <https://osm.etsi.org/>. Accessed on: Nov. 11, 2019.
- [11] R. Martínez, A. Mayoral, R. Vilalta, R. Casellas, R. Muñoz, S. Pachnicke, T. Szyrkowicz, and A. Autenrieth, "Integrated SDN/NFV orchestration for the dynamic deployment of mobile virtual backhaul networks over a multilayer (packet/optical) aggregation infrastructure", *Journal of Optical Communications and Networking*, vol. 9, no. 2, pp. A135-A142, Feb. 2017.
- [12] R. Casellas, R. Martínez, R. Vilalta, and R. Muñoz, "Control, Management, and Orchestration of Optical Networks: Evolution, Trends, and Challenges", *Journal of Lightwave Technology*, vol. 36, no. 7, pp. 1390–1402, Apr. 2018.
- [13] S. Fichera, R. Martínez, R. Casellas, B. Martini, R. Vilalta, R. Muñoz, and P. Castoldi, "Orchestrating Inter-DC Quality-Enabled VNFFG Services in Packet / Flexi-Grid Optical Networks," *European Conference on Optical Communication (ECOC)*, Rome, Italy, Sept. 2018.
- [14] S. Fichera, R. Martínez, B. Martini, M. Gharbaoui, R. Casellas, R. Vilalta, R. Muñoz, and P. Castoldi, "Latency-Aware Resource Orchestration in SDN-based Packet Over Optical Flexi-grid Transport Network," *Journal of Optical Communications and Networking*, vol.11, no. 4 pp. B83-B96, Apr. 2019.
- [15] A. Bravalheri, *et. al.*, "VNF Chaining across Multi-PoPs in OSM using Transport APL," *Optical Fiber Communication Conference (OFC)*, San Diego, USA, Mar. 2019.
- [16] F. J. Moreno-Muro, C. San-Nicolás, M. Garrich, P. Pavón-Mariño, O. González de Dios, and R. Lopez Da Silva, "Latency-Aware Optimization of Service Chain Allocation with Joint VNF Instantiation and SDN Metro Network Control," *European Conference on Optical Communication (ECOC)*, Rome, Italy, Sept. 2018.
- [17] M. Garrich, M. Hernández-Bastida, C. San-Nicolás, F.J. Moreno-Muro, P. Pavón-Mariño, "The Net2Plan-OpenStack Project: IT Resource Manager for Metropolitan SDN/NFV Ecosystems", *Optical Fiber Communication Conference (OFC)*, San Diego, USA, Mar. 2019.
- [18] J. Prados-Garzon, P. Ameigeiras, J.J. Ramos-Munoz, *et. al.* "Analytical Modeling for Virtualized Network Functions", in *Proc. of Inter. Conf. Comm. Work. (ICC Workshop)*, May 2017.
- [19] G. Faraci, A. Lombardo, "A simulative model of a 5G Telco Operator Network", *Procedia Comp. Sci.*, vol. 110, pp 344-351, 2017.
- [20] A. Alleg, T. Ahmed, M. Mosbah, *et. al.* "Delay-aware VNF placement and chaining based on a flexible resource allocation approach", in *Proc. of Int. Conf. Netw. Serv. Manag. (CNSM)*, Nov. 2017.
- [21] M. Savi, M. Tornatore, G. Verticale, "Impact of processing costs on service chain placement in network functions virtualization", in *Proc. of IEEE Conf. Netw. Func. Virt. Soft. Def. Netw. (NFV-SDN)*, Nov. 2015.
- [22] M. Gao, B. Addis, M. Bouet, S. Secci, "Optimal orchestration of virtual network functions", *Elsevier Com. Net.*, vol. 142, pp. 108-127, Sept. 2018.
- [23] D. Sattar, A. Mattraway, "Optimal Slice Allocation in 5G Core Networks", *Cornell University Library*, Feb. 2018.
- [24] L. Askari, A. Hmaity, F. Musumeci, and M. Tornatore "Virtual-network-function placement for dynamic service chaining in metro-area networks," in *proc. Optical Network Design and Modeling (ONDM)*, Dublin, Ireland, pp. 1-6, 2018.
- [25] J. Pedro, A. Eira, and N. Costa, "Metro Transport Architectures for Reliable and Ubiquitous Service Provisioning," in *proc. Asia Communications and Photonics Conference (ACP)*, Guangzhou, China, Nov. 2018.
- [26] J. Jay, "Low signal latency in optical fiber networks," *Proc. of the 60th IWCS Conference*, Charlotte, NC, USA, Nov. 2011.
- [27] Optelian white paper "A Sensible Low-Latency Strategy for Optical Transport Networks," July 2014. [Online]. Available: <https://www.optelian.com/wp-content/uploads/2016/10/Optelian-Low-Latency-Strategy-for-OTN-WP.pdf>
- [28] V. Bobrov, S. Spolit, and G. Ivanovs. "Latency causes and reduction in optical metro networks," *Optical Metro Networks and Short-Haul Systems VI*. Vol. 9008. International Society for Optics and Photonics, 2014.
- [29] Metro-Haul deliverable D2.3 "Network Architecture Definition, Design Methods and Performance Evaluation". [Online]. Available: <https://zenodo.org/record/3496939>. Accessed on: Nov. 11, 2019.
- [30] Governmental statistics, Murcia. [Online]. Available: [http://econet.carm.es/inicio/-/crem/sicrem/PU\\_padron/cifofl0/sec1\\_c1.html](http://econet.carm.es/inicio/-/crem/sicrem/PU_padron/cifofl0/sec1_c1.html). Accessed on: Nov. 11, 2019.
- [31] Governmental documents, Alicante. [Online]. Available: [http://documentacion.diputacionalicante.es/censo.asp?tfm\\_order=DESC&tfm\\_orderby=Pobl2017](http://documentacion.diputacionalicante.es/censo.asp?tfm_order=DESC&tfm_orderby=Pobl2017). Accessed on: Nov. 11, 2019.
- [32] P. Pavón-Mariño, "Optimization of Computer Networks: Modeling and Algorithms: A Hands-on Approach," Hoboken, NJ, USA: Wiley, 2016.
- [33] J. L. Romero-Gázquez, M. Garrich, F. J. Moreno-Muro, M. V. Bueno-Delgado, and P. Pavón-Mariño, "NIW: A Net2Plan-based Library for NFV over IP over WDM Networks," *Int. Conference on Transparent Optical Networks (ICTON)*, Angers (France), July 2019.
- [34] NIW Javadoc documentation. [Online]. Available: <http://www.net2plan.com/documentation/current/javadoc/api/com/net2plan/research/niw/networkModel/package-frame.html>. Accessed on: Nov. 11, 2019.
- [35] Net2Plan releases. [Online]. Available: <https://github.com/girtel/Net2Plan/releases>. Accessed on: Nov. 11, 2019.