

Machine Learning-Aided Resource Assignment in Space Division Multiplexed Elastic Optical Networks

Shrinivas Petale, Suresh Subramaniam

Department of Electrical and Computer Engineering, The George Washington University, Washington, DC, USA
{srpetale,suresh}@gwu.edu

Abstract—An effective solution to the looming capacity crunch is to use a fine-grained flexible frequency grid for elastic optical transmission and space division multiplexing (SDM) in conjunction with spectrally efficient modulation formats. The routing, modulation, core, and spectrum assignment (RMCSA) problem is an important lightpath resource assignment problem in SDM elastic optical networks (SDM-EONs). Intercore Crosstalk (XT) degrades the quality of parallel transmissions on different cores, and the RMCSA algorithm must ensure that XT constraints are met while maximizing network performance. There is a tradeoff between spectrum efficiency and XT tolerance - higher modulation formats are spectrum-efficient but are also less forgiving in terms of XT by allowing fewer connections to exist on adjacent cores on the overlapping spectrum. Many XT-aware RMCSA algorithms use a limit on the number of lit cores on the overlapping spectrum to ensure that XT constraints are satisfied. In this paper, we propose a machine learning (ML)-aided threshold optimization approach which improves the performance of RMCSA algorithms in the literature by up to three folds in terms of bandwidth blocking probability.

I. INTRODUCTION

Recent increases in bandwidth demands due to cloud-based services, 5G and 6G communications, high resolution game streaming, and data center networks can be fulfilled by space division multiplexed elastic optical networks (SDM-EONs) [1]. Advancements in coherent optical transmission have made it possible to fine-tune transmission parameters and increase spectral efficiency. SDM-EON enables parallel transmission of optical signals through multicore fibers (MCFs) with distance-adaptive multicarrier transmission [2], [3]. However, the quality of transmission (QoT) of signals transmitted through MCF degrades due to the intercore crosstalk (XT) between weakly coupled cores [4], [5].

In SDM-EON, lightpaths are routed through cores on the route's links with a set of contiguous and continuous frequency slices (FSs). Theoretically, a multifiber link and an MCF link bear the same capacity for transmission provided that the physical parameters and geometry of the cores are the same. However, the signal transmission in MCF does not resemble the one in multifiber link due to the compact structure of weakly coupled cores, and the QoT of a signal on a spectrum slice degrades due to XT. Therefore, it is important to find ways to address the effect of XT and handle it carefully [6]. The selection of multiefficient modulation formats (MFs) affects XT levels. Choosing a spectrally efficient MF leads

to spectrum saving but offers low tolerance for XT and shorter transmission reaches (TRs), and vice versa. Thus, the resource assignment problem becomes more complex with the selection of multiefficient modulation formats. Several RMCSA algorithms use a limit on the number of lit adjacent cores on overlapping spectrum to satisfy XT constraints [6], [7].

Recently, machine learning (ML) has been used to solve complex problems in optical networks [8]. The underlying relationship between the network features and output labels is learnt and then can be used to design network models. In this paper, we propose an ML-aided approach to learn such underlying relationship for better selection of MFs and achieves a good tradeoff between spectrum utilization and XT tolerance. A distinguishing feature of our approach is that it can be used with any RMCSA algorithm. In our recent work [9], we have developed a heuristic RMCSA algorithm for dynamic lightpath requests called Tridental Resource Assignment (TRA). In this paper, we apply our ML approach to several RMCSA algorithms including TRA, and demonstrate that significant further performance improvements are possible.

The paper is organized as follows. We review recent literature on the topic in Section II. The network model and problem statement are introduced in Section III. The ML-aided approach is presented in Section IV. Details of the TRA algorithm are presented in Sections V. Section VI presents simulation results, and the last section concludes the work.

II. RELATED WORK

The RMCSA problem involves the selection of resources, viz., route, MF, core and spectrum for incoming connection requests. Recent focus has been on the design of complex but efficient algorithms due to cheaper computing facilities [4], [10]. Initially, routing, spectrum, and core assignment (RSCA) algorithms for SDM-EON have been proposed in [11]–[13]. The RSCA problem is modeled and solved with the consideration of XT along with fragmentation, congestion, etc. The sorting of demand sizes, spectrum partitioning, and core prioritizing are used to reduce fragmentation and to avoid XT accumulation in a specific area of spectrum or network. Joint selection of core and spectrum to reduce XT and fragmentation for on-demand requirements is studied in [14]. The XT accumulation in a specific link can be avoided

with the help of multipath routing instead of single shortest path routing [15]. Multipath routing also helps reducing the bandwidth blocking probability (BBP) by offering alternate routes for assignment in a congested network.

Variants such as searching for earlier spectrum than earlier core index [13], [16] and considering maximum value, instead of average value, of XT experienced by each frequency slot (FS) in a desired set of FSs [16] are also studied. Various XT-based algorithmic solutions to reduce BBP are studied in [6] where instead of worst case estimations, a XT-aware approach with the dynamic calculation of XT levels is used to assign spectrum. The XT constraints can be used to model the effect of XT on signal transmission [6]. It is important to choose a significant XT threshold based on practical norms [17], which can be done by choosing a significant value of power coupling coefficient and span length [5], [7].

ML approaches have been used recently for specific problems such as QoT estimation [8] when there is scarcity of data [18], [19], using complex but efficient regression approaches [20], with data manipulation [19], [21], etc. However, there is a vast array of problems in optical networking for which the use of ML has not been sufficiently explored [4], [8], [22], [23]. A few studies such as [24] use unsupervised learning approach to solve resource allocation and fragmentation problems; however, supervised learning and reinforcement learning show better outcome [25]. For sub-problems of larger ones, static planning can be done using supervised learning to simplify dynamic provisioning [22], [26].

III. MODEL, PROBLEM STATEMENT, AND ILLUSTRATIVE EXAMPLE

We now present the network model and problem statement along with an example to illustrate how the XT thresholds affect the assignment of resources to connections.

A. Network Model and Problem Statement

We assume that the SDM-EON operates with a flexible grid of 12.5 GHz granularity and is equipped with coherent transceivers (TRXs). The TRXs support reconfigurable bit-rates and various MFs. The TRXs operate at a fixed baud rate of 14 GBaud, and each TRX transmits/receives an optical carrier allocated on 2 frequency slices (FSs) (i.e., 25 GHz). If the requested bit-rate exceeds the maximum capacity of a single TRX using a particular MF, the request is carried by several optical carriers within one superchannel (Sch). Each Sch is separated from neighbor Schs by 12.5 GHz guard-bands [16]. The nodes are connected using optical links consisting of MCFs. Each fiber link consists of the same set of MCFs with a specified core geometry in both the directions.¹ Two challenging and widely accepted core geometries, 3-core and 7-core [16], are studied. The effect of XT is dominant between the cores which are right next to each other called as *adjacent/neighbor cores*. For instance, each core in a 3-core

fiber has two adjacent cores. Similarly, in 7-core fiber, the outer cores have three adjacent cores and the center core has six adjacent cores (see Fig. 1). Furthermore, spatial continuity is imposed, which means that the same core is assigned to a lightpath on all MCF links on a route. Lightpath requests arrive at a specified Poisson rate with exponentially distributed mean holding time of unity (arbitrary units) and the datarates are uniformly distributed over a given range.

The set of modulation formats available for assignment is denoted as $D = \{f_1, f_2, \dots, f_{|D|}\}$, where f_1 is the lowest MF and $f_{|D|}$ is the highest MF. In this paper, we assume 5 modulation formats: QPSK (f_1), 8QAM, 16QAM, 32QAM, and 64QAM (f_5). The transmission reach (i.e., maximum length of a lightpath) depends on the lightpath's MF as well as the state of overlapping spectrum on adjacent cores (OsaCs).

The XT model of [27] is used to obtain the TRs for various MFs and for various values of lit adjacent cores (or litcores for short) for a XT^{th} of -40dB, as shown in Tab. I (from [28]). We consider average XT values of -40 dB between two adjacent cores after a single span of propagation [28]. The total XT experienced by a core is the sum of individual average XT contributions from each neighbor. For example, the TR for a 16QAM lightpath is 1950 km if the OsaC has a maximum of 0 litcores (i.e., the OsaC cannot be used for any other lightpath) and 900 km if the OsaC can be occupied by 6 lightpaths (i.e., all adjacent cores can be used by other lightpaths).

The problem is to assign resources (route, MF, core, and spectrum) to incoming connection requests so that the XT constraints are met for the incoming request and continue to be satisfied for all ongoing connections. The overall objective is to minimize the bandwidth blocking probability (BBP).

Incoming connection requests have a variety of choices of MFs that may satisfy the XT requirement. While higher MFs reduce the required spectrum, they also tighten up the core occupancy constraints for future connections, and a judicious balance between the two has to be struck. One approach to choosing an MF judiciously is to introduce a threshold for the number of allowable litcores, denoted γ^d for MF d . For instance, suppose $\gamma^3 = 4$. This means that if 16QAM (f_3) is to be assigned for a lightpath, then the lightpath's length cannot be larger than the TR for $\gamma^3 = 4$, i.e., 1100 km. Thus, if a request whose path length is 1200 km arrives, the setting of $\gamma^3 = 4$ means that 16QAM cannot be assigned to this request. Now, if γ^3 were set to be 3, then the TR becomes 1250 km, and the incoming request may be assigned 16QAM.

B. An Illustrative Example

We illustrate with an example the effect of the selection of thresholds on the tradeoff between spectrum utilization and the XT accumulation in the network with the help of Fig. 1 and Tab. I. The cross-section of the 7-core MCF is shown on the top to show the occupancy of cores in spatial domain and the corresponding occupancy of frequency slots is shown at the bottom. Three MFs, QPSK (f_1 , $d = 1$), 16QAM (f_3 , $d = 3$) and 64QAM (f_5 , $d = 5$), and their corresponding thresholds are used in this example. First fit policy is used in the selection of

¹In this paper, we assume that all the links have a single MCF in each direction, but the proposed work can be easily generalized for multiple fibers per link.

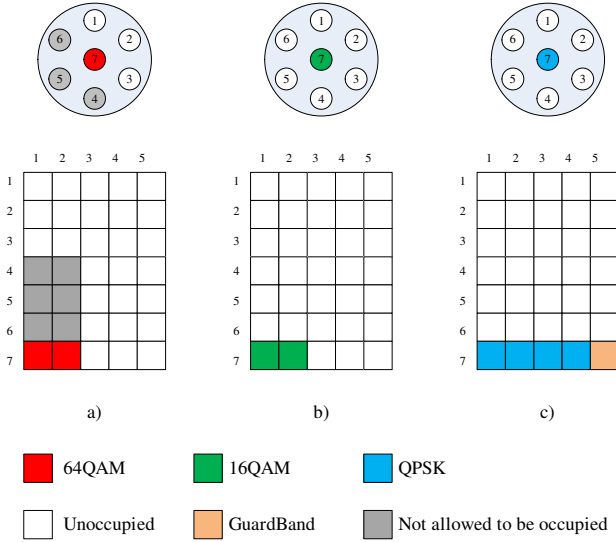


Figure 1: Change in spectrum utilization and XT levels for two different sets of γ^d thresholds.

MF. Suppose the first set of γ^d thresholds is $\{\gamma^1 = 5, \gamma^3 = 4, \gamma^5 = 3\}$ which corresponds to the TRs of 4750km, 1100km and 300km as shown in Tab. I. Suppose a connection request of 100Gbps arrives with a path length of 270km. The spectrum requirement for a 100Gbps connection using QPSK, 16QAM, and 64QAM as per the network model in Section III is 5 FSs (4 FSs and 1 Guardband), 2 FSs, and 2 FSs, respectively.

Starting from the highest MF, for 64QAM the TR for $\gamma^5 = 3$ is 300 km, and the TR values for larger values of γ are less than the path length. Therefore, if 64QAM is selected for this request, overlapping spectrum on any three cores out of six adjacent cores are allowed to be lit/occupied. In other words, overlapping spectrum on any three (total – allowed) cores is not allowed to be occupied as long as this connection exists in the network as shown in Fig. 1a. Thus the spectrum requirement effectively becomes 8 FSs (2 FSs + 3×2 FSs) if 64QAM is chosen. For 16QAM, the threshold is $\gamma^3 = 4$, and the corresponding TR is 1100 km. As the TR value of 900 km for $\gamma = 6$ is greater than the path length of 270 km, none of the adjacent cores are blocked by the current connection if 16QAM is chosen, which makes the spectrum requirement as 2 FSs as shown in Fig. 1b. Similarly, for QPSK, $\gamma^1 = 6$, and the spectrum requirement is 5 FSs as shown in Fig. 1c. Thus, the available choices of MFs for the connection are 64QAM, 16QAM, and QPSK.

Now, let us suppose that a different set of threshold values is chosen: $\gamma^1 = 5, \gamma^3 = 4$, and $\gamma^5 = 4$. Here, the TR value corresponding to $\gamma^5 = 4$ is only 250 km which is less than the path length, and therefore 64QAM is not a possible MF choice, and only QPSK and 16QAM are possible MF choices.

This example illustrates how setting the XT threshold values can affect the possible MF choices and thereby help the algorithm in finding a good balance between spectrum usage and XT tolerance. Note that setting all values of γ^d to 0 corresponds to always assigning the highest MF to a connection

at the expense of keeping all adjacent cores unoccupied. At the other extreme, setting all γ^d values to 6 corresponds to considering the worst-case XT scenario [6], [14], [16].

Table I: Transmission reaches (in km) of MFs for different values of allowable lit core (γ) for a 7-core MCF for a XT of -40dB per span of 50km (from Table II in [28]).

γ (Litcores)	Modulation Formats ($ D = 5, f_d \in D$)				
	QPSK	8 QAM	16 QAM	32 QAM	64 QAM
0	9050	3600	1950	1000	500
1	7650	3050	1650	850	400
2	6650	2650	1400	700	350
3	5850	2350	1250	650	300
4	5250	2100	1100	550	250
5	4750	1700	1000	500	250
6	4350	1500	900	450	200

IV. MACHINE LEARNING-AIDED LITCORE THRESHOLD SELECTION

We now present a novel machine learning-aided approach to select the litcore thresholds. This approach can be used by any RMCSA algorithm that uses litcore thresholds for meeting XT constraints. We discuss the prerequisites for the ML model (MLM) and its working principle. Recall that D is the set of MFs and γ^d is the litcore threshold for the d^{th} MF. The first step is to gather the samples of set of thresholds (STs) and the second step is to train the MLM. Let the optimal value of γ^d be denoted as γ_*^d . The steps in ML-aided Threshold Optimization (MLTO) to get γ_*^d for all MFs using MLM is shown in Fig. 2 and is explained below. The aim is to skip and change the MF using the γ^d threshold (SCT) using MLTO for reducing the BBP.

A. Data Aggregation/Approach of Measurements

Generating a single value of BBP from a ST needs a simulation of dynamic network operations. Thousands of such STs are possible from Tab. I, and each ST changes the network performance. The effect of selection of a particular ST can be seen by simulating that scenario. Since it takes a prohibitive amount of time to simulate the performance for all possible STs, we let the MLM learn the relationship between ST and BBP and let it select the optimal ST. A few hundreds of samples of STs are generated and are then used to obtain BBP values. Such STs are then fed to MLM which learns the relationship and then can predict the BBP of the remaining STs. The time required by MLM to train and predict is negligible as compared to the total computation time required for generating the BBP for each sample. In addition, the MLM can be easily modified by changing the hyperparameters, which makes the learning process even more flexible. The error in predicted and actual values of BBP of samples of STs with known BBPs is used to tune and calibrate the MLM for precise predictions.

In the first stage, several random STs are selected and the corresponding BBPs are generated using a given RMCSA algorithm. For each sample, the network model explained in Section III is used. A total of 100,000 connections arrive

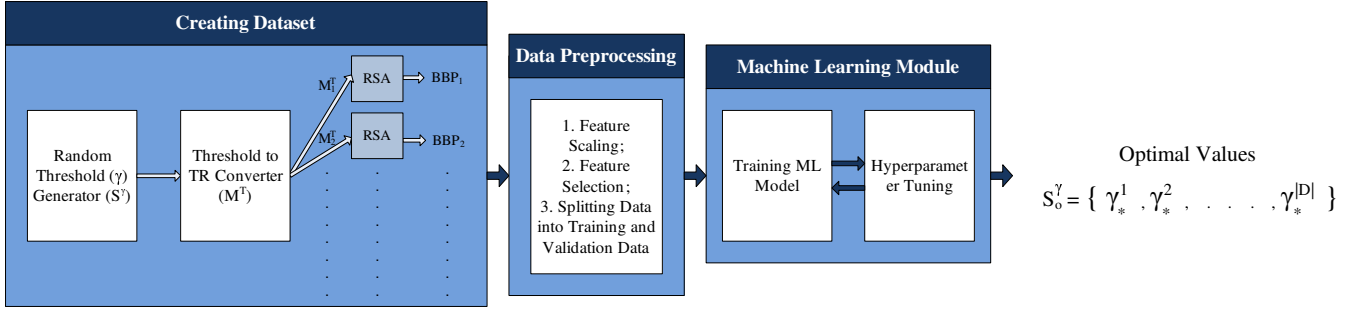


Figure 2: Machine learning approach for selection of γ_*^d .

dynamically and the network resources such as route, MF, core and spectrum are dynamically assigned based on the given RMCSA algorithm. The first 10,000 connections are discarded as warmup connection requests. Recall that the STs affect the selection of MFs changes which in turn impacts the spectrum utilization and XT accumulation, and thereby the BBP. Each γ^d in ST is selected such that the corresponding TR is nonzero,² to make sure that each MF is used. Let S_i^γ denote the i^{th} ST ($|S_i^\gamma| = |D|$). The STs are randomly generated. Each ST is used to generate the filtered TR model, denoted as M_i^T , from a TR model in Tab. I for a given baud rate and XT threshold per span. The filtered model is then used to dynamically select MFs based on path length of connection requests. It is important to consider Erlang loads that are high enough so that the different STs result in different BBPs, else, the BBPs for all STs may be 0, which is not useful for the MLM. The i^{th} sample for training MLM is comprised of $\{S_i^\gamma, BBP_i\}$ and is denoted as S_i .

B. Machine Learning-aided Threshold Selection

The STs represent features and the corresponding BBP values represent labels in S^γ for the MLM. The feature-label samples are then used for training the MLM. The MLM looks for statistical relationships and learns better if the dataset is preprocessed before it is fed to the MLM. The steps in training the model and optimizing the STs is given in Algo. 1.

Data preprocessing is the second step as shown in Fig. 2 and in Line 1 of Algo. 1. It involves feature selection, feature scaling, and separating samples into training and validation sets. Each feature value is checked for its contribution in the label. In other words, if a change in the feature does not affect the label then such feature is marked and discarded from the learning process. In our case, we observed that the change in each candidate γ^d threshold affects the BBP.

Now, in the final steps of data preprocessing the selected features are statistically processed. First the complete dataset is divided into training dataset and validation dataset (Line 2). In general, 80% of the samples are used for training and the remaining for validation. The features and labels in the training dataset are denoted as F_t and P_t , respectively. The

rest of the dataset is used for validation and denoted as F_v and P_v . Both F_t and F_v are scaled around the mean using the mean and standard deviation of F_t (Line 3). We use k-fold cross validation (kCV) to get the final assessment score to check the fit of the model. We use the inbuilt function of kCV in scikit-learn [29]. We use root mean squared error (RMSE) and R-squared (R2) score as assessment with inbuilt gridsearch cross validation approach to tune the hyperparameters (Line 4). Both of these assessments give similar results. The R2 score of 1.0 and RMSE of 0.0 means that the model has a perfect fit over the samples. As the MLM learns the smaller dataset and then is used to predict the BBP for rest of the STs, it has to learn the exact relationship between all STs and corresponding BBPs. The R2 score (or RMSE value) shows how well the MLM is fit to the dataset. It is desired to have the R2 score as 1.0 which means that the MLM has found a regression line which has a perfect fit to the dataset. However, the MLM struggles with overfitting which means that it fits perfectly to the known dataset and offers higher R2 score but predicts BBP with unacceptable error for unknown samples of STs and offers lower R2 score. Thus hyperparameter tuning using training and validation data is an essential step and it improves the performance of MLM by at least 28% in R2 score over default hyperparameters. The tuned hyperparameters are finally used to train the model where it learns the underlying relationship between STs and the corresponding BBP (Lines 6 and 7). The trained ML model using gridsearch procedure gives the set of γ_*^d denoted as S_o^γ (Line 8).

We observed that the relationship between the γ^d threshold values and BBP is not linear. We have tried various MLMs of regression such as Linear Regression, Polynomial Regression, Ridge Regression, Lasso Regression, ElasticNet Regression, Multi-layer Perceptron Regression, Support Vector Regression (SVR), etc., and observed that K-nearest Neighbour (KNN) regression model performs better in terms of increasing the R2 Score (or reducing RMSE) in the prediction phase and generating optimal outcomes when the training and validation model are kept the same for all the models. Thus in this paper, the results are presented for KNN MLM.

V. TRIDENTAL RESOURCE ASSIGNMENT

The MLM in the previous section can be combined with any RMCSA algorithm that uses XT thresholds. We will apply it

²A 0 value of TR (i.e., a TR that is shorter than the shortest network link) is possible in the case of XT of -25dB per span [28].

Algorithm 1 ML-Aided Optimization Model

Input: STs S_i^γ and corresponding BBP_i

Output: S_*^γ

- 1: Select the desired features
 - 2: Get F_t and F_v , and P_t and P_v from S^γ and BBP
 - 3: Get F_t^s and F_v^s by scaling F_t and F_v using the mean and standard deviation of F_t
 - 4: Choose hyperparameters by hyperparameter tuning using GridSearchCV()
 - 5: Feed type of ML model, set of hyperparameters and evaluation metric
 - 6: Get the desired hyperparameters using F_t and P_t , and F_v and P_v
 - 7: Train the model with obtained hyperparameters for which offers best value of the evaluation metric
 - 8: Use GridSearchCV() and trained model to get S_*^γ
-

to a FF RMCSA algorithm and another algorithm from the literature in Section VI. Recently, we proposed a novel RMCSA called Tridental Resource Assignment algorithm (TRA) in [9]. For reader convenience, we explain how the TRA algorithm works. We start with a few definitions and then present the algorithm.

We call a set of desired number of FSs as a slice window (SW). A different candidate SW is obtained by shifting the starting index of a current SW towards right by one on a given core. The *capacity*, *capacity loss (CL)*, and *tridental coefficient (TC)* of a candidate SW are defined as follows.

A. Capacity of an SW

The capacity of the n^{th} SW on k^{th} shortest path (SP) on route r , denoted as r^k , for a given MF f_d , denoted by $v_{n,c}^{k,d}$, is the number of cores on the whole path on which the SW can be assigned in the current network state (i.e., before resource assignment for the incoming request). Here, the current network state includes the litcore restriction of the already established connections on Overlapping spectrum on adjacent Cores (OsaCs). When an SW is assigned to the incoming request on a core with MF f_d , the capacity would decrease by an amount that depends on allowable γ^d of f_d and number of adjacent cores of the current core. Thus, the remaining capacity of the n^{th} SW on c^{th} core on r^k using f_d , denoted as $v_{n,c}'^{k,d}$, is the capacity of the SW if it were to be assigned to the incoming request, with the actual value of γ of selected MF f_d , denoted as γ_d , and the actual network state.

B. Capacity Loss and Tridental Coefficient of a Slice Window

The remaining capacity of a SW after the resource assignment of a connection varies based on the selected core and XT tolerance of the selected MF f_d . Thus, for every core and MF pair, $v_{n,c}'^{k,d}$ varies. The decrease in capacity after the hypothetical provisioning from the capacity before provisioning gives the total CL. Finally, the CL for the n^{th} SW on c^{th} core on r^k using f_d is calculated using (1).

$$\psi_{n,c}^{k,d} = v_{n,c}^{k,d} - v_{n,c}'^{k,d}. \quad (1)$$

The optimal choice of spectrum is when shared resources in the network are still available for future demands. When an SW on a path is assigned to a request, there is a CL for the SW on all the overlapping (shared) paths as well. Let, Z^k be the set of all the shared paths; $Z^k = i_1, i_2, \dots, i_z$. Let $\Delta(r, m)$ denote incoming request which arrived on route r and has datarate m . The number of FSs required to accommodate datarate m using MF f_d is denoted as β_d^m . We assume a lightpath tuple, $l_{\Delta(r,m)}(k, c, n, \beta_d^m)$, which represents the n^{th} SW of size β_d^m on c^{th} core on r^k for request $\Delta(r, m)$. The total CL of $l_{\Delta(r,m)}(k, c, n, \beta_d^m)$ is shown in (2); where, $\psi_{n,c}^{r^k,d}$ is the CL ($\psi_{n,c}^{k,d}$) on r^k and $\psi_{n,c}^{i_z,d}$ is the CL ($\psi_{n,c}^{k,d}$) on the z^{th} shared path i_z ($i_z \in Z^k$).

Finally, the TC of $l_{\Delta(r,m)}$ is defined as the sum of normalized values of CL, size of SW in terms of number of FSs, and the starting index of SW. It is denoted as $\Psi(l_{\Delta(r,m)})$ and is given in (4). The normalization is done using the respective maximum values, viz., maximum possible CL denoted by $\max \psi'(l_{\Delta(r,m)})$ given in (3), largest possible demandsize of datarate m denoted as β_1^m , and highest possible index of a SW is equal to $S - \beta_d^m + 1$ where S is the total number of FSs. The maximum CL can be C and thus $\max \psi'(l_{\Delta(r,m)})$ is obtained using (3).

$$\psi'(l_{\Delta(r,m)}) = \psi_{n,c}^{r^k,d} + \sum_{z=1}^{|Z^k|} \psi_{n,c}^{i_z,d}. \quad (2)$$

$$\max \psi'(l_{\Delta(r,m)}) = C + \sum_{z=1}^{|Z^k|} C = (1 + |Z^k|)C. \quad (3)$$

$$\Psi(l_{\Delta(r,m)}) = \frac{\psi'(l_{\Delta(r,m)})}{\max \psi'(l_{\Delta(r,m)})} + \frac{\beta_d^m}{\beta_1^m} + \frac{n}{S - \beta_d^m + 1}. \quad (4)$$

C. Tridental Resource Assignment (TRA)

We now briefly describe our proposed TRA algorithm. The algorithm we use here is a simplified version of the presented in [9].

As spatial continuity is imposed, a FS on a core is considered as free only if it is free on the same core index on all the links on the path. We refer the SW as *available* if this free SW can be occupied by current connection and does not affect ongoing connections on OsaCs. The pseudo-code for TRA is given in Algorithm 2.

All the possible SWs of spectrum for datarate m using MF f_d on all the cores are stored in set B_d^m . For datarate m , H^m is the set of all B_d^m sets for all MFs. The SW of index n on core c in B_d^m is denoted by $b_{c,n}^{m,d}$. For the datarate m , V^m is the set of β_d^m for all MFs. If the spectrum requirement of two MFs for a connection are the same, then the MF with larger spectrum is marked unavailable for the connection and such sorted MFs are stored in D^m . The TR of MF f_d for the corresponding value of γ is denoted as T_d^γ . The optimal

Algorithm 2 TRA Algorithm

Input: Network topology, $\Delta(r, m)$, set of SPs $P(r)$, their path lengths l_r^k , V^m , H^m

Output: $l_{\Delta(r, m)}^*(k^*, c^*, n^*, \beta^*)$

```

1:  $l_{\Delta(r, m)}^* \leftarrow \emptyset$ ,  $\Psi(l_{\Delta(r, m)}^*) \leftarrow \infty$ ,  $k \leftarrow 1$ 
2: while  $l_{\Delta(r, m)}^* = \emptyset \wedge k \neq K$  do
3:   for all ( $f_d \in D^m \wedge T_d^0 \geq l_r^k$ ) do
4:     lit core  $\gamma_d \leftarrow \gamma$  for which  $T_d^\gamma \geq l_r^k$ 
5:     Get  $\beta_d^m$  and  $\beta_1^m$  from  $V^m$ 
6:     for all  $b_{c, n}^{m, d} \in B_d^m$ ,  $B_d^m \in H^m$  do
7:       lightpath  $l_{\Delta(r, m)} \leftarrow (k, c, n, \beta_d^m)$ 
8:       if SW is available then
9:         Calculate  $\Psi(l_{\Delta(r, m)})$  of  $b_{c, n}^{m, d}$  SW
10:        if  $\Psi(l_{\Delta(r, m)}) < \Psi(l_{\Delta(r, m)}^*)$  then
11:           $l_{\Delta(r, m)}^* \leftarrow l_{\Delta(r, m)}$ ,  $k^* \leftarrow k$ ,  $c^* \leftarrow c$ ,  $n^* \leftarrow n$ ,  $\beta^* \leftarrow \beta_d^m$ 
12:        end if
13:      end if
14:    end for
15:  end for
16:  if  $l_{\Delta(r, m)}^* = \emptyset$  then
17:     $k \leftarrow k + 1$ 
18:  else
19:    break
20:  end if
21: end while

```

lightpath $l_{\Delta(r, m)}^*(k^*, c^*, n^*, \beta^*)$, corresponding TC and SP are initialized in Line 1. Here, the desired path index, core index, index of SW and demandsize are denoted by k^* , c^* , n^* , and β^* . In Line 2, the algorithm continues until either the SW with the lowest (best) TC is found or the search over all the SPs is completed. In Line 3, the search is initiated for those MFs in D^m whose maximum TR, i.e., TR without the consideration of XT at $\gamma = 0$ (T_d^0) is higher than path length of the k^{th} SP. In Line 4, the allowable litcore value γ_d is the γ value for which the TR value T_d^γ is greater than the path length of the k^{th} SP. In Line 5, the actual and largest sizes of SW are obtained. In Line 6, loop iterating over all the choices of SWs in B_d^m is started. The candidate lightpath l is initiated in Line 7. If the SW is available, the TC is calculated for the current SW $b_{c, n}^{m, d}$ on lightpath l in Line 9. In Lines 10-12, the information of SW which offers the least value of TC is stored as the desired lightpath l^* . In Line 17, the algorithm checks whether the desired lightpath $l_{\Delta(r, m)}^*$ is obtained or not. If it is obtained then the algorithm stops in Line 19; otherwise, the algorithm continues with the next SP in Line 17. Finally, after all the SWs on all the cores on the whole path are processed, the optimal lightpath $l_{\Delta(r, m)}^*$ is selected for a given connection, and the network resources are assigned to the connection request accordingly. If an SW is not found on all the SPs (i.e., $l_{\Delta(r, m)}^* = \emptyset$), the request is rejected.

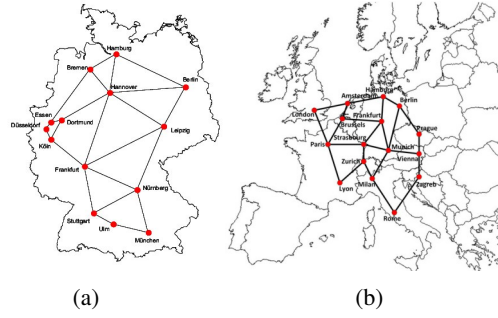


Figure 3: Network topologies a) DT (Fig. 3a), b) EURO (Fig. 3b).

VI. SIMULATION RESULTS

We now present simulation results comparing the chosen RSA with and without ML-aided optimization for a variety of scenarios. We use two practical topologies: generic German (DT) and European (EURO) shown in Fig. 3. The spectrum of 4 THz is considered on each link with each slice of 12.5 GHz ($\delta = 12.5$) i.e., 320 FSs ($S = 320$). Poisson connection arrival process with exponentially distributed holding time of 1 (arbitrary time unit) is assumed. The Erlang loads were chosen so that the BBP values generally range between 10^{-5} and 10^{-1} . A total of 100,000 requests are generated per trial, with the first 10,000 warm up requests being discarded. 95% confidence intervals with 10 trials are obtained for each experiment. The data rates are uniformly distributed between 40-400 Gbps with the granularity of 40 Gbps. There exist 3 SPs between every s-d pair ($K = 3$). Five MFs ($|D| = 5$), i.e., f_1 to f_5 are used viz., QPSK (f_1 , $d = 1$), 8 QAM, 16 QAM, 32 QAM and 64 QAM (f_5 , $d = 5$). The TR model for each MF with the average XT threshold per span of -40 dB with 14 GBaud TRX with the span length of 50 km is used from [28]. Around 400-1000 samples are generated and 80% of the samples are used for training, and the remaining for validation. The kCV uses $k=5$ folds to get assessment scores.

The improvement in performance of all the algorithms is observed for the average XT thresholds per span of -25dB [28], [30] and -40dB [28]; only the results for -40dB are presented due to space constraints. We compare the performance of algorithms such as a baseline XT-aware first fit (xtFF), P-XT in [16] and our algorithm TRA in [9] with and without the use of ML-aided optimization. xtFF chooses the highest possible MF and the first available slice window (SW) on the lowest numbered possible core for assignment. P-XT does XT-aware spectrum assignment with exhaustive search on all the routes. The spectrum available on the earliest index among the ones available on all the path-core pairs is selected. The algorithms are compared for the same parameters; especially, for the same XT^{th} , with imposed spatial continuity constraint, and spectrum choices available for spectrum search. K shortest paths are used in these algorithms. xtFF and TRA search for the next SP only if an SW is not available on the current SP. P-XT always searches the spectrum choices on all the SPs to choose the earliest SW on any core on all the SPs.

All the algorithms are XT-aware, meaning only SWs that satisfy the XT constraints are assigned. The new variants of the algorithms that use ML-aided Threshold Optimization (MLTO) are denoted with '-ML' in front of the algorithms' names.

The variation in BBP with respect to change in traffic loads in DT topology (Fig. 4a and Fig. 4b) and EURO topology (Fig. 4c and Fig. 4d) is shown in Fig. 4. The distribution of utilized MFs for all the algorithms with and without the use of MLTO for different values of C for XT^{th} of -40dB in DT topology (Fig. 5a and Fig. 5b) and EURO topology (Fig. 5c and Fig. 5d) is shown in Fig. 5. In all the scenarios of topology-core pairs, TRA outperforms P-XT and xtFF by a huge margin.³ TRA is designed to choose a MF while considering the effect of selection of MFs on the spectrum utilization-XT tradeoff. As explained in Section V, TRA calculates the loss in capacity for each MF-core pair by using the respective γ^d value for a specific SW. xtFF and P-XT do not give emphasis on the selection of MFs but on spectrum. We can observe from Fig. 5 that the distribution of MFs for TRA is quite different from xtFF and P-XT for all the scenarios. In addition, TRA shows a nearly uniform selection of all the MFs as compared to xtFF and P-XT which rely on choosing only 40-60% of the available MFs. It is also observed that the distribution of the MFs is similar for $C = 3$ and $C = 7$ for a given topology. Such minor variation for $C = 7$ is due to the presence of a central core which has six adjacent cores instead of a smaller and constant value (of three cores for non-central in $C = 7$ and two cores in $C = 3$).

The ML-aided version of each algorithm is obtained by using the γ_*^d thresholds generated by the MLM for SCT. These optimized thresholds are shown in Tab. II. The MLM learns the underlying and unquantifiable relationship between the samples of STs and corresponding BBP values and then generates the optimal set of thresholds. We observed that the MLM learns that the samples that do not follow the constraint are not essential and treated as outliers based on the corresponding BBP value. It is evident from the fact that the corresponding TRs of all the optimal outcomes by MLM follow the constraint in such cases. In addition, when the threshold generation is completely random, the ML automatically learns that keeping the threshold for the lowest MF (γ_*^1) such that the corresponding TR is high enough to accommodate longer paths is essential to lower connection blocking. Especially, when the TR at thresholds of higher MFs are lower, MLM learns the necessity of keeping γ_*^1 to low values to offer high TR to accommodate longer path lengths. However, we believe that this is due to the use of first fit policy in MF selection, and the results may vary if the first higher MF which can satisfy the QoT constraints is used. In both the policies, having constrained samples can always make the MLM learn efficiently and converge quickly. Setting up the

threshold of the lowest MF to a lower value to offer higher corresponding value of TR at threshold makes selection of MFs for longer paths easier and reduces the complexity of the MLM. However, we observed that the MLM learns the effect of TR of the lowest MF on BBP, and thus selects the γ_*^1 so that the corresponding TR is equal to or nearly equal to the highest value as shown in Tab. II.

Table II: ML-based S_*^γ for SCT.

Algorithm	Topology, C	Optimal Thresholds $\{\gamma_*^1, \gamma_*^2, \dots, \gamma_*^5\}$
TRA	DT, $C=3$	{0, 2, 2, 2, 0}
TRA	DT, $C=7$	{2, 6, 2, 3, 5}
TRA	EURO, $C=3$	{2, 2, 2, 1, 1}
TRA	EURO, $C=7$	{3, 3, 3, 3, 3}
P-XT	DT, $C=3$	{2, 2, 2, 2, 2}
P-XT	DT, $C=7$	{5, 5, 4, 5, 6}
P-XT	EURO, $C=3$	{1, 1, 2, 0, 2}
P-XT	EURO, $C=7$	{0, 6, 2, 0, 6}
xtFF	DT, $C=3$	{2, 2, 2, 2, 2}
xtFF	DT, $C=7$	{3, 3, 1, 3, 5}
xtFF	EURO, $C=3$	{2, 2, 2, 2, 2}
xtFF	EURO, $C=7$	{0, 3, 3, 5, 3}

It is important to know that each RSA algorithm reacts to a particular set of threshold values differently. We observed that the variation of BBP values with respect to the samples of STs is different for different algorithms for each scenario (not shown here). This is because STs decide the selection pattern of MFs for SCT. The selection approach is different for xtFF and P-XT, and TRA. The selection of spectrum choice is different for xtFF and P-XT, which in turn decides the future occupancy of remaining spectrum and thus affects the overall selection pattern of MFs. As the selection pattern is unique for each RSA algorithm, a common set of thresholds cannot be provided for a given topology-core scenario. Hence the STs are different for each scenario as shown in Tab. II.

The MLTO improves the selection pattern of MFs which results in reduced BBP values. In all the scenarios of topology-core pairs, the versions of RSA with the use of MLTO perform better than the actual algorithm without its use. However, the improvement is different for different algorithms for different scenarios. In case of $C = 3$, the performance of P-XT-ML is better than all the algorithms and their variants as shown in Fig. 4a and Fig. 4c. Here, the pattern of resource selection of P-XT-ML is similar to that of xtFF-ML but the selection of earliest spectrum choice reduces fragmentation to a greater extent. The modified selection pattern leverages this spectrum selection and improves the network performance. With the help of ML, the share of 16QAM is increased in the case of DT topology and the share of 8QAM is increased in the scenario of EURO topology for $C = 3$ as shown in Fig. 5a and Fig. 5c. This is because, the threshold values of 64QAM and 32QAM, and 16QAM are set to higher values which offer reduced TR values. However, the corresponding TR value for 16QAM is still higher than the TRs of 64QAM and 32QAM which increases its share. The increased selection of 16QAM, with reduction in the selection of 64QAM and 32QAM, saves spectrum and also offers the higher value of litcore (2, See Tab. II), which makes all the Overlapping spectrum

³Note that some BBP values are lower than 10^{-4} because some trials gave 0 blocked requests at loads. Confidence intervals have uneven lengths due to logarithmic scale on y-axis. The lengths are larger if the variation is huge for different traffic patterns.

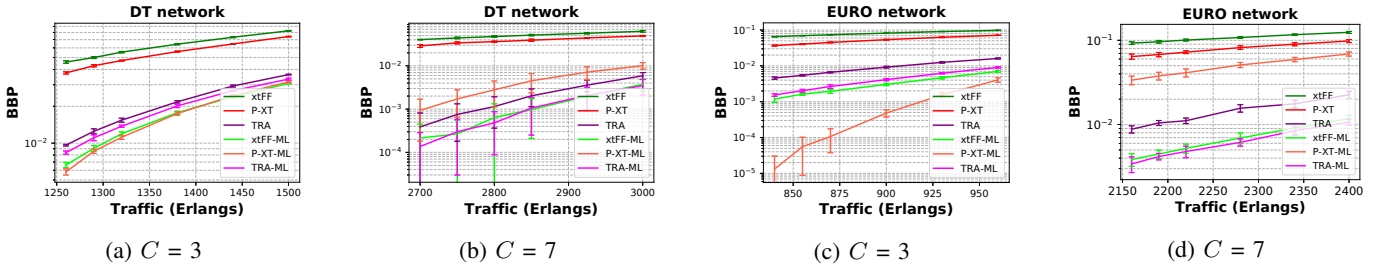


Figure 4: Variation in BBP wrt traffic for different values of C for XT^{th} of -40dB in DT (Fig. 4a and Fig. 4b) and EURO (Fig. 4c and Fig. 4d) topologies.

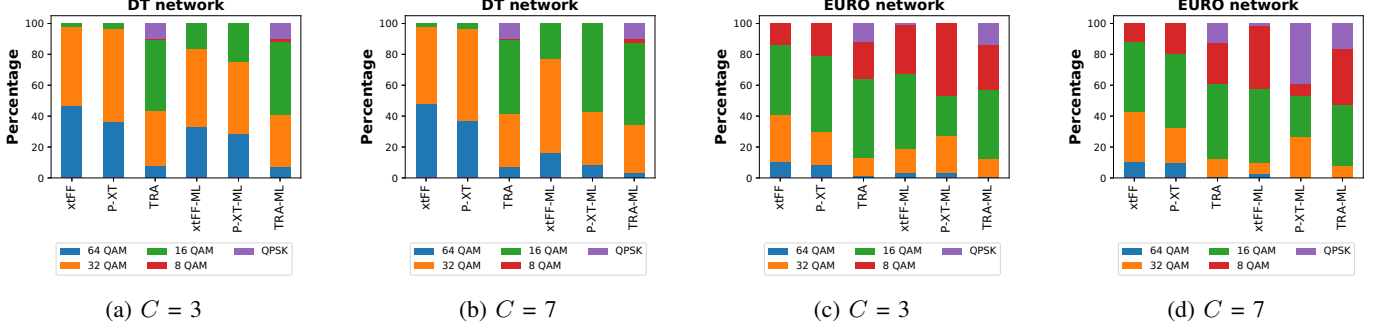


Figure 5: Distribution of utilized MFs for all the algorithms with and without MLTO for different values of C for XT^{th} of -40dB in DT (Fig. 5a and Fig. 5b) and EURO (Fig. 5c and Fig. 5d) topologies.

on adjacent Cores (OsaCs) available for future connections. A similar effect is observed for the EURO topology with higher selection of 8QAM after the use of MLTO. The optimal thresholds and distribution of MFs are different for both the topologies because of the difference in link lengths and node distributions which offer different path lengths of SPs in both the topologies. xtFF-ML performs better than TRA-ML because of better selection of MFs. A similar explanation for P-XT-ML is applicable to xtFF-ML as the selection pattern of MFs is similar for both the algorithms.

For the scenario of $C = 7$ with DT and EURO topologies, TRA-ML outperforms all the RSA algorithms and their variants as shown in Fig. 4b and Fig. 4d. The core geometry with $C = 7$ offers a different challenge of core selection as the selection of the central core has increased effect of XT on adjacent cores. TRA assigns the cores and MFs after considering the effect of the γ^d in terms of capacity loss and thus balances the tradeoff. Using MLTO, the selection of MFs improves a little more. Unlike the case of xtFF-ML and P-XT-ML in all scenarios, it is observed that TRA-ML shows an increased selection of (central MFs) 32QAM and 16QAM in DT topology (Fig. 5b), and 16QAM and 8QAM (Fig. 5d) as compared to TRA. The selection of central MFs offer higher values of γ^d and lower values of spectrum utilization. Proper selection of MFs manifests itself as reduced BBP, because when a more XT-sensitive MF is chosen it may save spectrum but prevents future connections to be assigned on OsaCs.

It is evident from the performance evaluation that TRA efficiently balances the tradeoff among various factors for

different core geometries and path lengths in different topologies. In addition, we can also observe that the variation in pattern of selection of MFs in RSA-ML is slightly different from the pattern in RSA; however, the decrease in the BBP is huge. Thus, a proper selection of MFs can improve the performance of the network dramatically, and ML can greatly help in selecting the optimal thresholds.

VII. CONCLUSION

We considered crosstalk-aware RMCSA algorithms for space division multiplexed elastic optical networks. There is a tradeoff between spectrum utilization and increased inter-core crosstalk (XT) depending on the selected modulation format. Many RMCSA algorithms use a constraint on the number of litcores on overlapping spectrum to satisfy XT constraints. In this paper, we proposed a machine learning-aided approach that optimizes the thresholds to control the selection of MFs. The approach can be used with any RMCSA allocation algorithm that limits the number of litcores to satisfy XT constraints. We compared the performance of a few RMCSA algorithms in the literature with the ML-aided variants, and showed that the ML-aided variants dramatically improve the bandwidth blocking probability of dynamic connection requests in a variety of scenarios. A very interesting observation is that the ML-aided variant of the worst performing baseline algorithm outperforms the best-performing algorithm without ML-aided threshold optimization.

Acknowledgement This work was supported in part by NSF grant CNS-1813617.

REFERENCES

- [1] B. Shariati, D. Klonidis, I. Tomkos, D. Marom, M. Blau, S. Ben-Ezra, M. Gerola, D. Siracusa, J. Macdonald, N. Psaila, *et al.*, “Realizing spectrally-spatially flexible optical networks,” *IEEE Photon, Society Newsletter*, pp. 4–9, 2017.
- [2] Y. Awaji, J. Sakaguchi, B. J. Puttnam, R. S. Luís, J. M. D. Mendinueta, W. Klaus, and N. Wada, “High-capacity transmission over multi-core fibers,” *Optical Fiber Technology*, vol. 35, pp. 100–107, 2017.
- [3] G. M. Saridis, D. Alexandropoulos, G. Zervas, and D. Simeonidou, “Survey and evaluation of space division multiplexing: From technologies to optical networks,” *IEEE Communications Surveys & Tutorials*, vol. 17, no. 4, pp. 2136–2156, 2015.
- [4] M. Klinkowski, P. Lechowicz, and K. Walkowiak, “Survey of resource allocation schemes and algorithms in spectrally-spatially flexible optical networking,” *Optical Switching and Networking*, vol. 27, pp. 58–78, 2018.
- [5] T. Hayashi, T. Taru, O. Shimakawa, T. Sasaki, and E. Sasaoka, “Design and fabrication of ultra-low crosstalk and low-loss multi-core fiber,” *Optics express*, vol. 19, no. 17, pp. 16576–16592, 2011.
- [6] M. Yang, Y. Zhang, and Q. Wu, “Routing, spectrum, and core assignment in sdm-eons with mcf: node-arc ilp/milp methods and an efficient x-aware heuristic algorithm,” *Journal of Optical Communications and Networking*, vol. 10, no. 3, pp. 195–208, 2018.
- [7] K. Walkowiak, A. Włodarczyk, and M. Klinkowski, “Effective worst-case crosstalk estimation for dynamic translucent sdm elastic optical networks,” in *ICC 2019-2019 IEEE International Conference on Communications (ICC)*, pp. 1–7, IEEE, 2019.
- [8] F. Musumeci, C. Rottondi, A. Nag, I. Macaluso, D. Zibar, M. Ruffini, and M. Tornatore, “An overview on application of machine learning techniques in optical networks,” *IEEE Communications Surveys & Tutorials*, vol. 21, no. 2, pp. 1383–1408, 2018.
- [9] S. Petale, J. Zhao, and S. Subramaniam, “Trident resource assignment algorithm for spectrally-spatially flexible optical networks,” in *ICC 2021-2021 IEEE International Conference on Communications (ICC)*, pp. 1–7, IEEE, 2021.
- [10] Í. Brasileiro, L. Costa, and A. Drummond, “A survey on challenges of spatial division multiplexing enabled elastic optical networks,” *Optical Switching and Networking*, p. 100584, 2020.
- [11] H. Tode and Y. Hirota, “Routing, spectrum and core assignment for space division multiplexing elastic optical networks,” in *2014 16th International Telecommunications Network Strategy and Planning Symposium (Networks)*, pp. 1–7, IEEE, 2014.
- [12] H. Tode and Y. Hirota, “Routing, spectrum and core assignment on sdm optical networks,” in *Optical Fiber Communication Conference*, pp. Tu2H–1, Optical Society of America, 2016.
- [13] H. Tode and Y. Hirota, “Routing, spectrum, and core and/or mode assignment on space-division multiplexing optical networks,” *Journal of Optical Communications and Networking*, vol. 9, no. 1, pp. A99–A113, 2017.
- [14] S. Fujii, Y. Hirota, H. Tode, and K. Murakami, “On-demand spectrum and core allocation for reducing crosstalk in multicore fibers in elastic optical networks,” *Journal of Optical Communications and Networking*, vol. 6, no. 12, pp. 1059–1071, 2014.
- [15] H. M. Oliveira and N. L. da Fonseca, “Multipath routing, spectrum and core allocation in protected sdm elastic optical networks,” in *2019 IEEE Global Communications Conference (GLOBECOM)*, pp. 1–6, IEEE, 2019.
- [16] M. Klinkowski and G. Zalewski, “Dynamic crosstalk-aware lightpath provisioning in spectrally-spatially flexible optical networks,” *IEEE/OSA Journal of Optical Communications and Networking*, vol. 11, no. 5, pp. 213–225, 2019.
- [17] L. Zhang, N. Ansari, and A. Khreishah, “Anycast planning in space division multiplexing elastic optical networks with multi-core fibers,” *IEEE Communications Letters*, vol. 20, no. 10, pp. 1983–1986, 2016.
- [18] D. Azzimonti, C. Rottondi, A. Giusti, M. Tornatore, and A. Bianco, “Active vs transfer learning approaches for qot estimation with small training datasets,” in *Optical Fiber Communication Conference*, pp. M4E–1, Optical Society of America, 2020.
- [19] C. Rottondi, R. di Marino, M. Nava, A. Giusti, and A. Bianco, “On the benefits of domain adaptation techniques for quality of transmission estimation in optical networks,” *Journal of Optical Communications and Networking*, vol. 13, no. 1, pp. A34–A43, 2021.
- [20] M. Ibrahim, H. Abdollahi, C. Rottondi, A. Giusti, A. Ferrari, V. Curri, and M. Tornatore, “Machine learning regression for qot estimation of unestablished lightpaths,” *Journal of Optical Communications and Networking*, vol. 13, no. 4, pp. B92–B101, 2021.
- [21] T. Panayiotou, G. Savva, B. Shariati, I. Tomkos, and G. Ellinas, “Machine learning for qot estimation of unseen optical network states,” in *Optical Fiber Communication Conference*, pp. Tu2E–2, Optical Society of America, 2019.
- [22] Y. Zhang, J. Xin, X. Li, and S. Huang, “Overview on routing and resource allocation based machine learning in optical networks,” *Optical Fiber Technology*, vol. 60, p. 102355, 2020.
- [23] J. A. Gordon, A. Battou, M. P. Majurski, D. Kilper, U. Celine, M. Tornatore, J. Pedro, J. Simsarian, J. Westdorp, and D. Zibar, “Summary: workshop on machine learning for optical communication systems,” 2020.
- [24] S. Trindade and N. L. da Fonseca, “Machine learning for spectrum defragmentation in space-division multiplexing elastic optical networks,” *IEEE Network*, 2020.
- [25] X. Chen, B. Li, R. Proietti, H. Lu, Z. Zhu, and S. B. Yoo, “Deepprmsa: A deep reinforcement learning framework for routing, modulation and spectrum assignment in elastic optical networks,” *Journal of Lightwave Technology*, vol. 37, no. 16, pp. 4155–4163, 2019.
- [26] R. J. Pandya, “Machine learning-oriented resource allocation in c+ l+ s bands extended sdm-eons,” *IET Communications*, vol. 14, no. 12, pp. 1957–1967, 2020.
- [27] P. Poggiolini, “The gn model of non-linear propagation in uncompensated coherent optical systems,” *Journal of Lightwave Technology*, vol. 30, no. 24, pp. 3857–3879, 2012.
- [28] C. Rottondi, P. Martelli, P. Boffi, L. Barletta, and M. Tornatore, “Crosstalk-aware core and spectrum assignment in a multicore optical link with flexible grid,” *IEEE Transactions on Communications*, vol. 67, no. 3, pp. 2144–2156, 2018.
- [29] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay, “Scikit-learn: Machine learning in Python,” *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.
- [30] J. Sakaguchi, B. J. Puttnam, W. Klaus, Y. Awaji, N. Wada, A. Kanno, T. Kawanishi, K. Imamura, H. Inaba, K. Mukasa, *et al.*, “305 tb/s space division multiplexed transmission using homogeneous 19-core fiber,” *Journal of Lightwave Technology*, vol. 31, no. 4, pp. 554–562, 2012.