



UNIVERSIDADE FEDERAL DA BAHIA

DISSERTAÇÃO DE MESTRADO

**ESTRATÉGIAS DE PLANEJAMENTO DE REDES ÓPTICAS
ELÁSTICAS CONSIDERANDO TRÁFEGO MULTIPERÍODO**

LEONARDO ALMEIDA JACOBINA MESQUITA

Programa de Pós-Graduação em Engenharia Elétrica

Salvador
9 de agosto de 2019

PPGE-Msc-2019

LEONARDO ALMEIDA JACOBINA MESQUITA

**ESTRATÉGIAS DE PLANEJAMENTO DE REDES ÓPTICAS
ELÁSTICAS CONSIDERANDO TRÁFEGO MULTIPERÍODO**

Esta Dissertação de Mestrado foi apresentada ao Programa de Pós-Graduação em Engenharia Elétrica da Universidade Federal da Bahia, como requisito parcial para obtenção do grau de Mestre em Engenharia Elétrica.

Orientador: Dr. Karcus Day Rosário de Assis

Salvador
9 de agosto de 2019

Ficha catalográfica elaborada pelo Sistema Universitário de Bibliotecas (SIBI/UFBA),
com os dados fornecidos pelo(a) autor(a).

Mesquita, Leonardo Almeida Jacobina,
Estratégias de Planejamento de Redes Ópticas
Elásticas Considerando Tráfego Multiperíodo /
Mesquita, Leonardo Almeida Jacobina. -- Salvador,
2019.

63 f. : il

Orientador: Assis, Karcus Day Rosário.
Dissertação (Mestrado - Engenharia Elétrica) --
Universidade Federal da Bahia, Escola Politécnica,
2019.

1. Redes de Computadores. 2. Comunicações Ópticas.
3. Programação Linear. 4. Alocação de Espectro. 5.
Planejamento Multiperíodo. I. Assis, Karcus Day
Rosário, . II. Título.

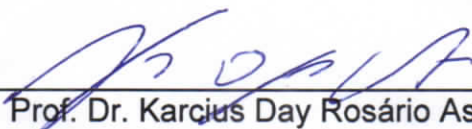
Leonardo Almeida Jacobina Mesquita

"Estratégias de Planejamento de Redes Ópticas Elásticas Considerando Tráfego Multiperíodo"

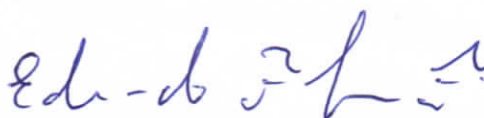
Dissertação apresentada à Universidade Federal da Bahia, como parte das exigências do Programa de Pós-Graduação em Engenharia Elétrica, para a obtenção do título de *Mestre*.

APROVADA em: 09 de Agosto de 2019.

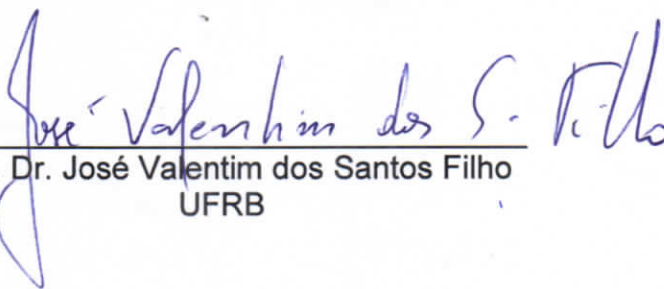
BANCA EXAMINADORA



Prof. Dr. Karcus Day Rosário Assis
Orientador - UFBA



Prof. Dr. Eduardo Furtado de Simas Filho
UFBA



Prof. Dr. José Valentim dos Santos Filho
UFRB

*Não podemos viver apenas para nós mesmos. Mil fibras nos conectam
com outras pessoas.*

—HERMAN MELVILLE

RESUMO

Os avanços na tecnologia de transmissão óptica de dados permitiram o grande crescimento da Internet e de serviços que exigem muita largura de banda. Simultaneamente, veio também a crescente demanda por maior velocidade de conexão por parte dos usuários. O paradigma mais recente para aumentar a eficiência das redes *backbone* de comunicação é composto pelas redes ópticas elásticas ou *Spectrum-Sliced Optical Networks* (SLICE) baseadas em multiplexação por divisão ortogonal de frequência ou *Orthogonal Frequency-Division Multiplexing* (OFDM). As redes baseadas em OFDM são mais eficientes que as antecessoras, baseadas em multiplexação por divisão de comprimento de onda ou *Wavelength Division Multiplexing* (WDM), no entanto o roteamento é mais complexo, pois deve levar em consideração a modulação, continuidade e contiguidade de espectro; originando assim o problema de roteamento, modulação e alocação de espectro ou *Routing, Modulation and Spectrum Assignment* (RMSA). A resolução do RMSA é fundamental para melhor uso dos recursos espectrais da rede óptica, já que a otimização desse problema é capaz de disponibilizar mais banda para demandas futuras. Considerando que o tráfego das redes *backbone* está em constante crescimento, é possível também assumir uma abordagem incremental de tráfego e resolver o RMSA para várias fases de chegada de tráfego. Este trabalho propõe uma nova formulação e algoritmos capazes de resolver o RMSA para períodos sequenciais de tráfego incremental. Uma nova formulação de programação linear foi proposta para otimizar a o espectro máximo utilizado nos enlaces da rede, mas devido ao fato do RMSA ser um problema cujo tempo para chegar a uma solução é dado como polinomial não-determinístico (NP), a resolução exige um longo tempo processamento quando o problema envolve um grande número de instâncias. Este é o caso de redes de grande dimensão. Diante disso, também foram propostas duas heurísticas para resolver o RMSA incremental em um tempo viável. Uma de roteamento por caminho mais curto e outra de roteamento por balanceamento de carga. Ambas as heurísticas foram simuladas com tráfego multiperíodo incremental, gerado estocasticamente, para estudar os efeitos a longo prazo de utilizar tais heurísticas em relação aos recursos espectrais. Também são propostas redes neurais artificiais para fazer previsão de tráfego com o objetivo de fazer testes com valores realistas de tráfego, gerados por redes neurais. Para confirmar a viabilidade da previsão, dois modelos propostos são avaliados. Tanto as simulações com tráfego aleatório quanto as simulações com tráfego proveniente de previsão mostraram padrões similares e forneceram resultados que são fundamentais para fazer o planejamento de redes ópticas com vistas a atender um tráfego futuro e ao mesmo tempo economizar recursos.

Palavras-chave: Redes de Computadores, Comunicações Ópticas, SLICE, Redes Elásticas, Algoritmos, Formulação, Programação Linear, Roteamento, Modulação, Alocação

de Espectro, Tráfego, Incremental, Planejamento, Multiperíodo, Redes Neurais Recorrentes, Previsão de Tráfego

ABSTRACT

The advances in optical data transmission technology have allowed the huge growth of the Internet and bandwidth demanding services. Alongside this phenomenon, came the growing demand for faster speed connection and even more available bandwidth. The most recent paradigm to increase the efficiency of communication systems is composed by the spectrum-sliced optical networks (SLICE) based on orthogonal frequency-division multiplexing (OFDM). OFDM-based networks are more efficient than its predecessor, based on wavelength division multiplexing (WDM), although the routing is more complex because it has to consider spectrum continuity, contiguity and modulation; thus originating the routing, modulation and spectrum assignment (RMSA) problem. The resolution of the RMSA is one of the keys to increase the spectral efficiency in SLICE networks, since the optimization of this problem is capable of providing more bandwidth for future traffic in a network. Considering that the traffic within *backbone* networks is constantly growing, it's possible to assume an incremental traffic approach and solve the RMSA for several phases of incoming traffic. This work proposes new formulations and algorithms capable of solving the RMSA for sequential periods of incremental traffic. A new linear programming formulation was proposed to optimize the maximum utilized spectrum in the network's links, but due to the fact that the RMSA is a NP problem, the resolution requires long processing time when the problem involves a large number of instances. This is the case for networks of larger dimensions. In face of that, heuristics were also proposed to solve the incremental RMSA in viable time. Two heuristics were proposed to solve the incremental RMSA. A shortest path heuristic and a load balancing heuristic. Both heuristics were simulated with stochastically generated multiperiod incremental traffic, in order to study the long-term effects that using such heuristics can have relative to the spectral resources. Artificial neural networks are also proposed to make traffic prediction with the objective of realizing tests with realistic traffic values, generated by the neural networks. To confirm the viability of neural networks in making traffic forecasting, two models are evaluated. Both the simulations with random traffic and the simulations with forecasted traffic have presented similar patterns and have provided results which are useful for network operators to make incremental planning of optical networks.

Keywords: Computer Networks, Optical Communication, SLICE, Elastic Networks, Algorithms, Formulation, Linear Programming, Routing, Modulation, Spectrum Assignment, Incremental Traffic, Planning, Multiperiod, Recurrent Neural Network, Traffic Prediction

SUMÁRIO

Capítulo 1—Introdução	3
1.1 Introdução	3
1.2 Trabalhos Relacionados	5
1.2.1 Planejamento de Redes e Tráfego Incremental	5
1.2.2 Previsão de Tráfego	6
1.3 Contribuições desta Obra	6
1.3.1 Artigos Publicados pelos Autores Durante a Pesquisa	7
Capítulo 2—Roteamento, Modulação e Alocação de Espectro	9
2.1 Multiplexação por Divisão Ortogonal de Frequência	9
2.1.1 Formatos de Modulação	10
2.1.2 Multiplexadores ópticos	12
2.1.3 Simulação de Formatos de Modulação	14
2.2 O problema do RMSA	14
2.3 Tráfego Incremental	16
Capítulo 3—Métodos de Resolução do RMSA Incremental	19
3.1 Formulação de Programação Linear	19
3.1.1 Notações	19
3.1.2 Parâmetros de Entrada	20
3.1.3 Variáveis	20
3.1.4 Descrição Matemática da Formulação	21
3.2 Heurísticas de RMSA Incremental	24
3.2.1 Variáveis e Funções	25
3.2.2 Algoritmo Base	26
3.2.3 Caminho Mais Curto com Modulação para Carga Mínima (SPMLM)	26
3.2.4 Balanceamento e Modulação para Carga Mínima (BMLM)	27
3.3 Resultados do RMSA Incremental	28
3.3.1 Redes Pequenas	29
3.3.2 Redes Grandes	30
3.3.3 SPMLM vs BMLM	31

Capítulo 4—Previsão de Tráfego	35
4.1 Redes Neurais Artificiais	35
4.1.1 Redes Neurais Feed-Foward	35
4.1.2 Redes Neurais Recorrentes	37
4.1.3 Long Short-Term Memory	38
4.2 Previsão de Matrizes de Tráfego	39
4.2.1 Banco de Dados	39
4.2.2 Modelos de RNNs com LSTM	39
4.3 Treinamento	41
4.4 Resultados do RMSA Incremental	42
Capítulo 5—Conclusões	45
5.1 Discussão e Conclusões Finais	45
5.2 Trabalhos Futuros	46

LISTA DE FIGURAS

2.1	Disposição de oito subportadoras ortogonais de banda estreita que somadas formam um único canal banda larga. Note que o valor de pico de cada subportadora ocorre quando a amplitude das outras portadoras é nula.	11
2.2	Diagrama de constelação para 16QAM. Os pontos azuis representam cada um dos 16 símbolos. Logo acima de cada ponto, se encontra a sequência de 4 <i>bits</i> equivalente	11
2.3	Diagrama do funcionamento de um OADM. O OADM demultiplexa o sinal recebido, redireciona alguns comprimentos de onda para outra parte, adiciona novos comprimentos de onda, provenientes de um sinal diferente e os multiplexa novamente para a saída.	12
2.4	Exemplo de resolução do RMSA para uma rede anel de quatro enlaces.	15
2.5	Resultado da simulação para duas matrizes de tráfego incremental, relativas a dois períodos incrementais consecutivos, $t = 1$, representado em azul e $t = 2$ representado em laranja.	18
3.1	Rede pequena com seis nós, usada para as simulações nesta secção	29
3.2	Número de SFs consumidos (Carga) por enlace da rede óptica Européia por quatro períodos de tráfego incremental. (a) SPMLM e (b) BMLM.	30
3.3	Diagrama de caixa da probabilidade de bloqueio relativo a cada período de tráfego incremental simulado. Apenas um formato de modulação foi usado, tanto para o SPMLM quanto para o BMLM.	33
3.4	Diagrama de caixa da probabilidade de bloqueio relativo a cada período de tráfego incremental simulado. Quatro formatos de modulação foram simulados para o SPMLM e o BMLM.	34
3.5	Cuvas da probabilidade de bloqueio média para os 25 períodos incrementais, e diferentes conjecturas de formatos de modulação.	34
4.1	Ilustração de um neurônio artificial.	36
4.2	Rede Neural Artificial Feed-Foward	37
4.3	Rede Neural Recorrente como uma junção sequencial de FFNNs convencionais	38
4.4	Diagrama da memória longa de curto prazo (LSTM)	39
4.5	Arquitetura da RNN proposta	40
4.6	Série temporal original e prevista para a demanda de tráfego do um par de nós anônimos origem-destino da rede GEANT.	42
4.7	Série temporal original e prevista para a demanda de tráfego do um par de nós anônimos origem-destino da rede Abilene.	42

4.8	Probabilidade de bloqueio para o período de tráfego incremental, para ambas as heurísticas de RMSA, com dados gerados pela RNN para a rede Abilene.	43
4.9	Probabilidade de bloqueio para o período de tráfego incremental, para ambas as heurísticas de RMSA, com dados gerados pela RNN para a rede GEANT	44

LISTA DE TABELAS

2.1	Formatos de Modulação.	14
3.1	Simulações com MILP	29
3.2	Simulações com Heurísticas	31
3.3	Parâmetros de redes e simulações. A referências com os parâmetros reais das redes: http://www.av.it.pt/anp/on/refnet2.html	32
4.1	Precisão das RNNs para as previsões de tráfego das redes Abilene e GEANT	43

LISTA DE ACRÔNIMOS

API	<i>Application Programming Interface</i>
BMLM	<i>Balanced Minimum Load Modulation</i>
BPSK	<i>Binary Phase Shift Keying</i>
EoL	<i>End-of-Life</i>
FFNN	<i>Feed-Foward Neural Network</i>
IBM	<i>International Business Machines Corporation</i>
LSTM	<i>Long Short-term Memory</i>
MILP	<i>Mixed-Integer Linear Programming</i>
MSE	<i>Mean Square Error</i>
NP	<i>Non-Deterministic Polynomial time</i>
OADM	<i>Optical Add-Drop Multiplexer</i>
OEO	<i>Óptico-Elétrico-Óptico</i>
OFDM	<i>Orthogonal Frequency-Division Multiplexing</i>
OXC	<i>Optical Crossconnect</i>
QAM	<i>Quadrature Amplitude Modulation</i>
QoS	<i>Quality of Service</i>
QPSK	<i>Quadrature Phase Shift Keying</i>
RMSA	<i>Routing, Modulation and Spectrum Assignment</i>
RNA	<i>Redes Neurais Artificiais</i>
RNN	<i>Recurrent Neural Networks</i>
ROADM	<i>Reconfigurable Optical Add-Drop Multiplexer</i>
RPROP	<i>Resilient Backpropagation</i>
RSA	<i>Routing and Spectrum Assignment</i>
SF	<i>Slot de Frequência</i>
SLE	<i>Static Lightpath Establishment</i>
SLICE	<i>Spectrum-Sliced Optical Networks</i>
SPMLM	<i>Shortest Path with Minimum Load Modulation</i>
WDM	<i>Wavelength Division Multiplexing</i>
WSS	<i>Wavelength Selective Switches</i>

INTRODUÇÃO

1.1 INTRODUÇÃO

Os últimos anos têm conhecido um grande aumento em novos serviços na internet, desde novas plataformas de entretenimento e redes sociais a transações financeiras, criptomoedas e computação em nuvem. Esses serviços são um fator comum para a maioria da população nos tempos atuais e a demanda por eles só tende a aumentar. Pesquisas realizadas pela Cisco mostraram que a taxa de dados transmitida apresentará um crescimento composto de 26% ao ano entre 2017 e 2022 (CISCO, 2017). Tais serviços são exigentes em banda, mas o acesso a *internet* de alta velocidade é amplamente difundido, graças aos avanços na tecnologia de fabricação de fibra óptica (AGRAWAL, 2013).

No entanto, pesquisadores têm enfrentado dificuldades em extrair mais capacidade das fibras ópticas monomodo, devido ao efeito Kerr, onde a potência de transmissão é limitada devido ao aumento do índice de refração na fibra de sílica ser proporcional à intensidade óptica (RAYMOND, 1999). Dado este fenômeno, a eficiência espectral ainda pode ser aumentada pelo uso de formatos de modulação de alta ordem (WINZER, 2012) e de diferentes técnicas de multiplexação, das quais a mais promissora é a multiplexação por divisão ortogonal de frequência ou *Orthogonal Frequency-Division Multiplexing* (OFDM) (LOWERY; DU; ARMSTRONG, 2007).

Este trabalho foca no caso especial de redes ópticas transparentes baseadas em OFDM, chamadas de *Spectrum-Sliced Optical Networks* (SLICE). As redes SLICE são mais flexíveis e possuem maior granularidade espectral de suas portadoras unitárias, também conhecidas como *slots* de frequência (SF). O tamanho atual desses *slots* de frequência variam entre 25 GHz, 12,5 GHz e 6,5 GHz (VELASCO et al., 2012). Estas propostas de tamanho de SF são mais adequados para o tráfego de dados atuais, que costuma ser bastante diverso. Em comparação, sistemas anteriores, baseados em multiplexação por divisão de comprimento de onda ou *Wavelength Division Multiplexing* (WDM), possuíam SF de tamanho igual a 50 GHz, segundo o padrão da ITU-T (JONCIC; HAUPT; FISCHER, 2012a).

Junto com a maior granularidade das redes SLICE também existe uma maior complexidade associada ao roteamento de tráfego, já que para rotear as demandas de tráfego, o algoritmo deve considerar não só os caminhos a serem percorridos, mas também quais serão os SF atribuídos a cada uma das conexões de forma a atender certas condições de alocação espectral. O problema de roteamento, somado às condições que devem ser atendidas, definem o problema conhecido como o problema de roteamento, modulação e alocação de espectro ou *Routing Modulation and Spectrum Assignment* (RMSA) (PIÓRO; MEDHI, 2004) apresentado no Capítulo 2.

O RMSA pode ser resolvido com o objetivo de otimizar algum parâmetro da rede, como, por exemplo, o número de *slots* consumidos, carga máxima por enlace ou fragmentação. Esses métodos podem ser encontrados nos trabalhos de Yin, Gong e Takagi (YIN et al., 2013; GONG et al., 2012; TAKAGI et al., 2011).

O RMSA pode ser avaliado de uma perspectiva estática, em que toda a demanda a ser oferecida para a rede, é conhecida e alocada de vez. No entanto, o RMSA também pode ser abordado de uma forma mais realística, assumindo uma perspectiva multiperíodo, na qual novas demandas surgem e desaparecem ao longo do tempo. Dentro desta abordagem, há o caso especial de considerar o tráfego como incremental, onde novas demandas surgem para serem alocadas, mas nunca desaparecem. Este último caso é o mais apropriado para algoritmos direcionados a planejar a rede para uso a longo prazo e para relatórios sobre carregamento máximo da rede (pior caso de consumo de recursos) (SOBH et al., 2008).

Publicações recentes (VELASCO et al., 2016; SOUMPLIS et al., 2018; MORALES et al., 2017; IYER; SINGH, 2018) focaram no planejamento de redes usando algoritmos específicos para tráfego incremental, mas nenhuma formulação de programação linear completa, que resolva o RMSA para tráfego incremental, foi proposta na literatura científica. Para sanar este problema, no Capítulo 3, é proposta uma formulação de programação linear inteira mista ou *Mixed-Integer Linear Programming* (MILP).

Como o tempo necessário para se obter resultados, na resolução do modelo de MILP, cresce exponencialmente com o número de variáveis introduzidas, é inviável usar estes modelos para simulações com redes grandes, como é o caso da maioria das redes continentais de *backbone*. Por este motivo, no Capítulo 3, também são propostas heurísticas capazes de resolver o RMSA incremental em um tempo consideravelmente menor.

O objetivo é fazer simulações com a formulação e as heurísticas para encontrar os valores estimados de capacidade exigida e probabilidade de bloqueio para novas conexões, que são métricas que irão aumentar ao longo do tempo, à medida que novas conexões irão ocupando os recursos disponíveis nos enlaces da rede. Esses dados podem ser usados para fazer o planejamento de redes (YOO; QIAO; DIXIT, 2000) e evitar custos adicionais como sobredimensionamento ou subdimensionamento da camada física.

Em geral, o planejamento de redes a partir da resolução do RMSA pode ser usado para prever custos operacionais e prover melhor qualidade de serviço ou *Quality of Service* (QoS), mas esse planejamento é mais útil quando são usados dados reais do tráfego na rede. (AZZOUNI; PUJOLLE, 2017a) mostrou que previsão de tráfego pode ajudar com planejamento de capacidade a longo termo, fornecendo meios de projetar redes do zero sem ter que depender da análise estática *End-of-Life* (EoL), onde resultados já mostraram que é menos eficiente, em comparação com o planejamento incremental de

tráfego (SOUMPLIS et al., 2018).

Neste trabalho, para as simulações, são usados dados de demanda de tráfego gerados de forma estocástica, para as redes onde o histórico real não estava disponível. No entanto, para fazer simulações mais realistas, foi usado, também, o histórico proveniente de um banco de dados contendo os valores de tráfego instantâneo. As simulações com histórico de tráfego foram feitas apenas para a rede Abilene e GÉANT, as quais possuem um banco de dados publicamente disponível (UHLIG et al., 2006).

Adicionalmente, para aumentar o período de simulação e testar os métodos de previsão de tráfego, foi usada uma rede neural recorrente ou *Recurrent Neural network* (RNN) com *Long Short-term Memory* (LSTM) com o intuito de fazer uma regressão a partir dos dados de tráfego existentes. Tanto a RNN quanto o modelo para previsão de tráfego são aprofundados no Capítulo 4.

As RNNs já foram usadas anteriormente em (AZZOUNI; PUJOLLE, 2017a) e (ZHAO et al., 2018a) para fazer previsão de matriz de tráfego, mas estes trabalhos individuais focaram em apenas uma rede individual com um intervalo de treinamento curto, para suas arquiteturas de RNN. Neste trabalho, novas arquiteturas de RNN são propostas. Os modelos de RNN foram, treinados e empregados para fazer a regressão que teve como saída uma matriz de séries temporais. Estas matrizes, por sua vez foram amostradas em intervalos constantes e convertidas em matrizes de tráfego incremental para que a simulação com as heurísticas fosse possível.

Após as simulações de RMSA incremental das redes Abilene e GEANT usando matrizes de tráfego providas do banco de dados junto com as matrizes de tráfego geradas pelas RNNs, os resultados são discutidos no Capítulo 4 e as conclusões finais da obra são apresentados no Capítulo 5. Trabalhos futuros e complementares são encontrados no Capítulo 5.2.

1.2 TRABALHOS RELACIONADOS

A recente mudança do paradigma do modelo de banda fixa do WDM para as redes com banda flexível (OFDM) abriu muitas possibilidades para a melhoria das redes ópticas a partir de melhores algoritmos de escalonamento ou de técnicas de resolução do problema de roteamento e alocação de espectro. As redes ópticas elásticas receberam o foco das pesquisas, mas este campo não é tão bem explorado como o das redes baseadas em WDM, no que diz respeito à análise multiperíodo e incremental de tráfego. Alguns trabalhos relacionados são apresentados neste capítulo.

1.2.1 Planejamento de Redes e Tráfego Incremental

O estudo feito por Straub (Straub; Kirstadter; Schupke, 2006) teve resultados onde o planejamento incremental de tráfego foi superior ao planejamento estático baseados na abordagem de fim de vida (*End-of-LIFE*) para arquiteturas WDM. O planejamento incremental apresentou valores quase ótimos nos casos em que a previsão de tráfego não era possível, ou se fazia pouco confiável. Também, para redes WDM, métodos computacionais para otimização forem providos por Lardeux (Lardeux; Nace, 2007), para operações

com tráfego incremental.

Eira (Eira; Pedro; Pires, 2015) apresenta uma comparação dos padrões de banda fixa (WDM) e banda flexível (OFDM) para dispositivos ópticos como *transceivers* e *line cards*, empregando um método de planejamento multiperíodo. Soumplis (SOUMPLIS et al., 2018) apresenta um algoritmo para planejar redes para tráfego incremental considerando também a degradação progressiva no extrato físico das fibras, indicando melhorias de até 36% para as redes de banda flexível. Nós também referimos o leitor ao algoritmo de Gkamas (GKAMAS; CHRISTODOULPOULOS; VARVARIGOS, 2015) capaz de resolver o RMSA e o roteamento da camada IP usando tráfego incremental, presumindo incremento uniforme ao longo do tempo. O algoritmo anterior foi usado por Iyer (IYER; SINGH, 2018), junto a uma formulação de ILP para cálculo de gastos, para fazer o planejamento de redes com uma abordagem multiperíodo incremental.

Em (VELASCO et al., 2016), os autores consideram um algoritmo de planejamento incremental da capacidade para tráfego *on-Demand* capaz de diminuir a probabilidade de bloqueio para valores próximos de zero, com a instalação de catorze placas de linha por mês.

1.2.2 Previsão de Tráfego

A habilidade de prever o estado futuro de uma rede se faz desejável para aplicações atuais como escalonamento de *slots* de frequência, detecção de DDoS e prevenção contra congestionamento, que se faz particularmente importante para serviços de computação em nuvem (AIBIN, 2018).

Métodos de regressão como ARIMA, FARIMA, previsão por base de *wavelet* e redes neurais do tipo *Feed-Foward* (FFNN) já foram comparados para fazer previsão de tráfego em (FENG; SHU, 2005) e (MOAYEDI; MASNADI-SHIRAZI, 2008).

Eventualmente, o foco mudou para previsão de matrizes de tráfego e trabalhos como (ZHAO et al., 2018a) e (AZZOUNI; PUJOLLE, 2017a) propuseram arquiteturas de redes neurais recorrentes com memória LSTM para fazer previsão de matrizes de tráfego para as redes Abilene e GEANT, respectivamente. A precisão obtida da regressão por RNN foi maior do que obtida por outras técnicas e é considerado estado da arte para dada tarefa.

Também, (MORALES et al., 2017) apresenta uma arquitetura de gerenciamento de redes OFDM baseada em previsão de dados. Nesta arquitetura, os modelos preditivos são usados para reconfigurar a rede de forma a minimizar os custos totais da rede. Esta produção conseguiu diminuir de oito a quarenta e dois por cento do número de transponders instalados.

1.3 CONTRIBUIÇÕES DESTA OBRA

A maior parte dos trabalhos apresentados na Secção 1.2 trata apenas de soluções para resolver o RSA, RMSA, RMSA incremental ou previsão de tráfego, de forma independente. Nenhum modelo de MILP para resolver o RMSA com tráfego incremental foi proposto, sem a ajuda de algoritmos e heurísticas. Para contribuir com a pesquisa neste campo,

no Capítulo 3, é proposta uma formulação de MILP para sanar este problema de forma otimizada, sem a necessidade de heurísticas.

Além do modelo de MILP para resolução do RMSA, nenhum estudo foi feito considerando o efeito a longo prazo de se usar heurísticas de RMSA. Esse estudo é importante, pois o esgotamento o consumo de recursos na rede pode ser diferente, a depender da heurística. Este estudo é feito no Capítulo 3 com uma análise da capacidade máxima exigida das fibras e da probabilidade de bloqueio para os respectivos algoritmos propostos. Os dados incrementais de capacidade e probabilidade de bloqueio podem ser usados para fazer planejamento de redes e estimativa de custos com equipamentos.

Esta dissertação também aborda a viabilidade da previsão de matrizes de tráfego e propõe uma arquitetura de RNN com LSTM para a rede Abilene e uma para a rede GEANT no Capítulo 4. Em seguida, é apresentada uma técnica para transformar as matrizes de fluxo de tráfego instantâneo em matrizes de tráfego incremental. A lista de contribuições pode ser declarada como:

- Proposta de formulação de MILP para resolver o RMSA incremental
- Proposta de heurísticas de RMSA incremental
- Propostas de arquiteturas de RNNs com LSTM para previsão de matrizes de tráfego
- Relatório incremental da capacidade exigida e da probabilidade de bloqueio em várias redes para planejamento multiperíodo

Em resumo, este trabalho visa fornecer métodos para o dimensionamento de redes ópticas elásticas a partir de simulações multiperíodo que fornecem estimativas de capacidade de máxima exigida pelas fibras e probabilidade de bloqueio na rede.

1.3.1 Artigos Publicados pelos Autores Durante a Pesquisa

Mesquita, A. J. L., & Assis, K.D.R., Santos, A.F. & Alencar, M. S., and Almeida Jr, R. (2018). A Routing and Spectrum Assignment Heuristic for Elastic Optical Networks under Incremental Traffic. 1-5. 10.1109/SBFoton-IOPC.2018.8610937.

Mesquita, L. A. J., Assis, K. D. R. Análise sobre redes Ópticas com Heurísticas de Roteamento, Modulação e Alocação de Espectro para Tráfego Incremental. In: The Brazilian Symposium on Computer Networks and Distributed Systems (SBRC). [S.l.: s.n.], 2019.

Assis, K.D.R., Mesquita, L. A. J., Almeida Jr, R. C. ; Waldman, Hélio ; Hammad A.; SIMEONIDOU D. Virtualization of optical networks utilizing squeezing protection: The exact formulation. Electronics Letter, p. 1-3, 2019.

ROTEAMENTO, MODULAÇÃO E ALOCAÇÃO DE ESPECTRO

Este capítulo trata do problema de roteamento, modulação e alocação de espectro, definido para redes ópticas elásticas, mas antes, se faz necessário uma introdução sobre temas intermediários. A Secção 2.1 fala sobre multiplexação por divisão ortogonal de frequências. A Secção 2.1.1 fala sobre formatos de modulação e a Secção 2.1.2 aborda multiplexadores ópticos para redes SLICE. Com esses conceitos abordados, a Secção 2.1.3 mostra dados e explica como os formatos de modulação são simulados na Secção 2.2, que dá um exemplo do problema do RMSA.

2.1 MULTIPLEXAÇÃO POR DIVISÃO ORTOGONAL DE FREQUÊNCIA

Multiplexação por divisão ortogonal de frequência (OFDM) é um esquema de sinalização para comunicação digital e o seu desenvolvimento teve começo em 1950 (MOSIER; CLA-BAUGH, 1958). A conceitualização de multiplexação por divisão de frequência foi feita por (CHANG, 1966) e (SALTZBERG, 1967). O propósito principal era fazer transmissão simultânea de dados com canais sobrepostos de forma a aumentar a eficiência espectral (ZOU; WU, 1995). No início da tecnologia, era extremamente caro e complexo implementar o OFDM devido ao grande número de osciladores e filtros necessários (ZOU; WU, 1995), mas com os avanços na tecnologia de integração em larga escala, ficou mais acessível a aquisição de *chips* de alta velocidade para calcular a transformada rápida de Fourier. Hoje, redes baseadas em OFDM já são populares.

Em redes OFDM, os símbolos do sinal a ser transmitido são modulados por portadoras igualmente espaçadas, de forma que as frequências das portadoras são ortogonais e o sinal a ser transmitido pode ser representado como:

$$v(t) = \sum_{k=0}^{N-1} X_k e^{j2\pi kt/T}, \quad 0 \leq t \leq T \quad (2.1)$$

onde X_k representa o k -ésimo símbolo, N é o número de subportadoras e T é a duração do símbolo. Como k é um número inteiro, as frequências ortogonais são harmônicas da frequência dada por $1/T$. Logo, as portadoras estarão espaçadas de forma regular em frequências múltiplas inteiras de $1/T$.

Sejam duas frequências $2\pi k_1/T$ e $2\pi k_2/T$, a ortogonalidade pode ser provada, caso a integral da Equação (2.2) seja verdadeira, que obviamente é o caso.

$$\frac{1}{T} \int_0^T (e^{j2\pi k_1 t/T})^* (e^{j2\pi k_2 t/T}) dt = 0. \quad (2.2)$$

Assumimos que a o espectro de frequência de cada símbolo seja representado pela função *sinc* definida pela Equação (2.3).

$$\mathcal{F}\{X_k\} \implies V(f) = A \frac{\sin(\pi T f)}{\pi f} \quad (2.3)$$

Em redes OFDM, se cada X_k a ser transmitido em paralelo for modulado por subportadoras de frequências ortogonais, então, no domínio da frequência, vários sinais *sinc* estarão sobrepostos e estreitamente separados em um intervalo de $1/T$. A Figura 2.1 mostra o arranjo de oito bandas de frequência estreitamente espaçadas. Durante a transmissão, as bandas estreitas são somadas para formar um canal de banda larga. Note que na frequência central equivalente a qualquer uma das portadoras, a amplitude das demais portadoras é igual a zero, sendo assim, possível distinguir as subportadoras individuais na recepção do sinal. É importante lembrar que as subportadoras, são associadas com comprimentos de onda e transmitidos, em sistemas ópticos.

2.1.1 Formatos de Modulação

As redes SLICE empregam transmissores flexíveis com suporte para diferentes níveis de modulação (ZHAO et al., 2014), onde o número de *bits* por símbolo pode ser determinado por formatos de modulação como *Binary Phase-Shift Keying* (BPSK), *Quadrature Phase Shift Keying* (QPSK) ou *Quadrature Amplitude Modulation* (QAM).

Esquemas de modulação como BPSK, QPSK e QAM representam uma sequência de *bits* como um símbolo codificado em forma de um segmento senoidal. Esses símbolos são diferenciados pela amplitude e/ou fase. A Figura 2.2 mostra uma constelação 16QAM. O formato de modulação QAM usa 16 símbolos, ilustrados na constelação por 16 pontos (Q, I). Cada ponto define um fasor que é usado para representar a sequência de $\log_2 16 = 4$ *bits* que é transmitida como um segmento de senoide com mesmo ângulo e amplitude do fasor.

16QAM usa 16 símbolos. Já o 8QAM usa 8 símbolos. O número de *bits* transmitidos por símbolo é igual a $\log_2 M$, onde M é o número de símbolos, geralmente indicado no prefixo da nomenclatura do formato de modulação. QAM diferencia seus símbolos usando fasores com diferentes ângulos e amplitudes. QPSK é análogo, mas a amplitude é sempre constante e os símbolos são diferenciados apenas pelo ângulo. Já o BPSK usa apenas dois fasores possíveis para representar os seus símbolos.

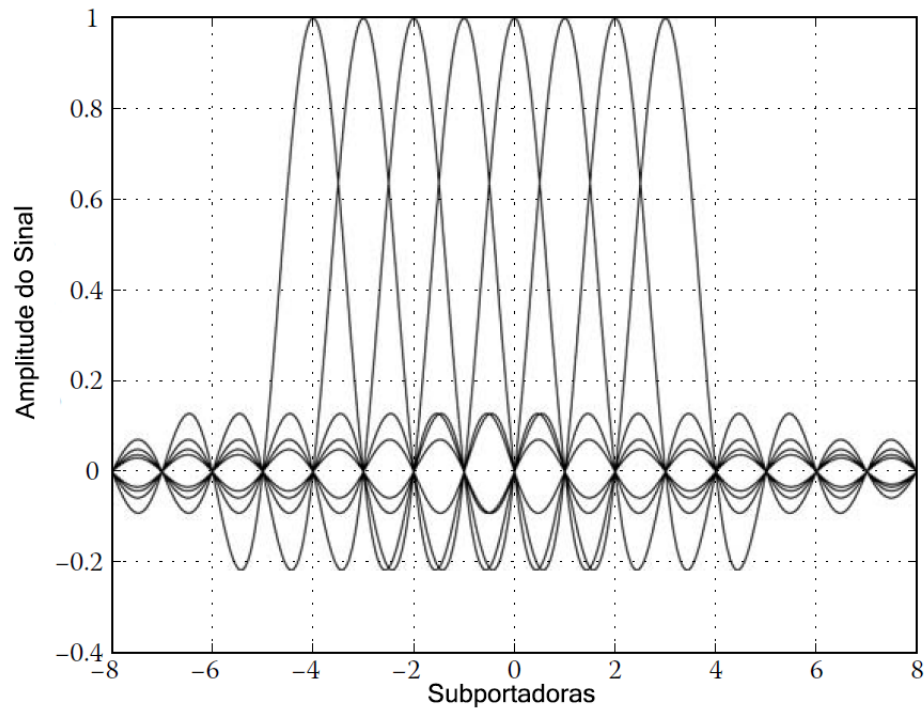


Figura 2.1: Disposição de oito subportadoras ortogonais de banda estreita que somadas formam um único canal banda larga. Note que o valor de pico de cada subportadora ocorre quando a amplitude das outras portadoras é nula.

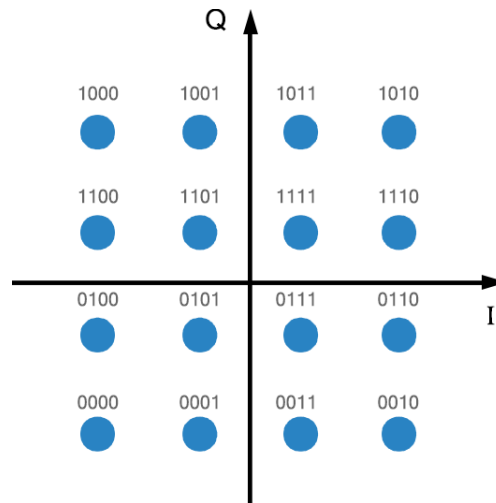


Figura 2.2: Diagrama de constelação para 16QAM. Os pontos azuis representam cada um dos 16 símbolos. Logo acima de cada ponto, se encontra a sequência de 4 *bits* equivalente

Taxas de transmissão mais altas são possíveis com níveis de modulação mais altos, mas, em sistemas de comunicações ópticas a razão sinal-ruído é significativamente baixa nos receptores, limitando o alcance de transmissão e criando a necessidade de lasers mais potentes. Apesar disso, constelações de 1024QAM e até mesmo 2048QAM já estão

em estado avançado de pesquisa e foram até demonstrados funcionando (BEPPU et al., 2015).

2.1.2 Multiplexadores ópticos

Para transmitir os sinais modulados e suas respectivas subportadoras, é preciso fazer uso de *switches* ópticos nos nós presentes entre as fibras ópticas. No início, os *switches* eram eletrônicos e a comutação dependia de conversão óptica-elétrica-óptica (OEO), o que consumia muita potência e aumentava a latência de transmissão.

Eventualmente, *switches* fotônicos e tecnologias ópticas foram usadas para construir dispositivos capazes de fazer a comutação e seleção de comprimentos de onda sem precisar de conversão OEO. Essa seleção de comprimentos de onda é feito por intermédio de dispositivos conhecidos como *Optical Add-Drop Multiplexer* (OADM), ilustrado na Figura 2.3. Em um ponto de intersecção da rede, após demultiplexar os sinais sobrepostos, os OADMs podem escolher alguns comprimentos de onda para serem roteados por uma saída e os restantes por outra saída. A função de remover comprimentos de onda é dado como "Drop". Os OADMs também podem inserir novos comprimentos de onda carregando sinais de outra origem na saída desejada. Essa função é dada como "Add".

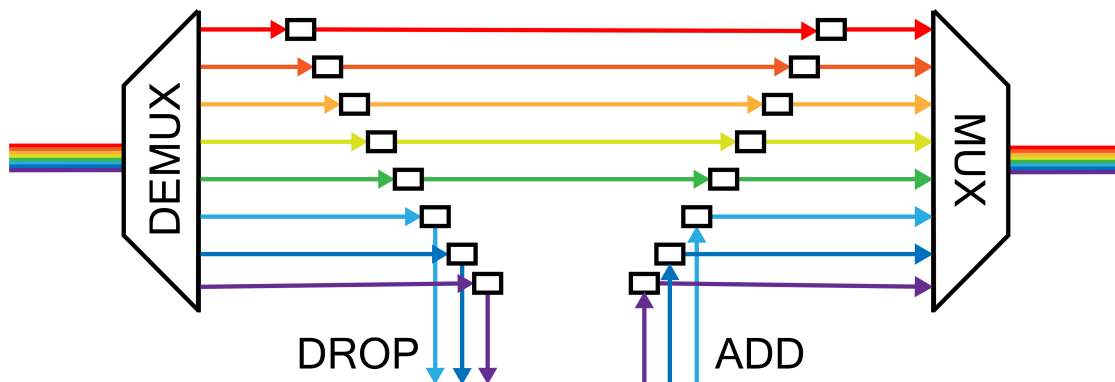


Figura 2.3: Diagrama do funcionamento de um OADM. O OADM demultiplexa o sinal recebido, redireciona alguns comprimentos de onda para outra parte, adiciona novos comprimentos de onda, provenientes de um sinal diferente e os multiplexa novamente para a saída.

As funções do OADM são executadas de forma completamente óptica, diminuindo a energia consumida pelos *switches*, diminuindo a latência de processamento nos nós e tornando a ampliação da rede menos custosa (ADAMS, 2008). Existem dois tipos de OADM que podem ser evidenciados: o primeiro tipo é o que usa *Wavelength Selective Switches* (WSS), sendo essa uma técnica de direcionamento de comprimento de onda por intermédio de microespelhos localizados dentro do dispositivo. O segundo tipo de OADM é o que usa uma matriz de conectores ópticos que leva cada comprimento de onda de cada uma das entradas para cada uma das saídas, conhecidos como *Optical Crossconnects* (OXCs) (ELDADA, 2005).

Algumas arquiteturas de OADM são reconfiguráveis e a seleção pode ser feita de forma remota. Esses dispositivos são chamados de *reconfigurable optical add-drop multiplexer*.

(ROADM) (ROORDA; COLLINGS, 2008).

Redes ópticas baseadas em OFDM, tais como as redes SLICE, possuem uma divisão mais granular de seu espectro de frequência, implicando em um maior número de subportadoras de frequência, comumente chamados de *slots* de frequência (SF). A banda selecionada para cada subportadora pode ser tão pequena quanto 12.5 GHz (CHRISTODOULOPOULOS; TOMKOS; VARVARIGOS, 2011), em comparação com redes baseadas em WDM, cuja banda fixa de frequência por portadora costuma ser 50 GHz nos padrões mais recentes como o da ITU-T (JONCIC; HAUPT; FISCHER, 2012b). Essa é uma grande vantagem para lidar com diferentes demandas de tráfego pois reduz os desperdícios com banda não utilizada na faixa da subportadora, aumentando a eficiência espectral da rede (DJORDJEVIC; VASIC, 2006).

Em redes completamente ópticas, como a SLICE, os dados são roteados ao longo dos enlaces, sem precisar passar por nenhum tipo de interface óptico-elétrico-óptico (OEO). Por isso, não é possível fazer conversão de espectro. Isso significa que, quando uma conexão for atendida pela rede, a ela será associado um número de SFs, de acordo com o taxa de transmissão e o formato de modulação usado. Esses SFs, então, definem uma banda de frequência específica que será ocupada pela conexão e, como não haverá conversão de espectro, a faixa espectral consumida deve ser a mesma em todas as fibras em que a conexão deve passar durante a etapa de roteamento. Esta peculiaridade das redes completamente ópticas define a restrição conhecida como restrição da continuidade, que deve ser considerada nos algoritmos que serão apresentados posteriormente.

Outra restrição imposta sobre às redes SLICE, por questões de eficiência, força que os SFs pertencentes a uma mesma conexões sejam alocados adjacentes (i.e. contigualmente) uns aos outros, evitando que sejam intercalados por banda de guarda, diminuindo o número de subportadoras consumidas, em geral. Esta regra define a condição conhecida como restrição da contiguidade. Quando o roteamento das conexões é feito considerando essas restrições, o problema é nomeado de problema de roteamento e alocação de espectro.

Somado a estes fatores, a possibilidade de usar transmissores/receptores adaptáveis que permitem usar múltiplos formatos de modulação, espectros de largura de banda e taxas de transmissão, dá às redes SLICES uma maior eficiência espectral (JINNO et al., 2010; ZHAO et al., 2014), tornando-as a melhor opção, no momento, para atender a crescente demanda por transmissão de dados. Se mais de um nível de modulação pode ser usado, então também terá que ser decidido qual a melhor opção, considerando a banda consumida e o alcance de transferência de cada formato de modulação. Essas considerações tornarão o processamento mais complexo e caracterizam o problema conhecido como roteamento, modulação e alocação de espectro (RMSA).

Existem outras regras que devem ser obedecidas, para evitar sobreposição de espectros referentes a diferentes conexões, manter a conservação de fluxo de dados, respeitar os limites de capacidade da fibra, entre outros, que estão implícitos no problema e que devem ser considerados nos algoritmos de resolução do RMSA.

2.1.3 Simulação de Formatos de Modulação

Os ROADMs disponíveis no mercado, assim como os *transceivers* adaptáveis, permitem a escolha do número de bits por símbolo para processamento no *transponder* (CHRISTODOULOPOULOS; TOMKOS; VARVARIGOS, 2011) e isso é considerado durante o RMSA. A escolha de Formatos de modulação torna a tarefa mais complexa já que os formatos de modulação mais eficientes possuem menor alcance de transmissão. A Tabela 2.1 mostra os formatos de modulação que foram julgados realistas para as distâncias envolvidas nas redes de *backbone* simuladas posteriormente. Os valores da tabela estão de acordo com publicações recentes (OLIVEIRA; FONSECA, 2017; JAYA; GOPINATHAN; RAJENDRAN, 2016) e o valor de banda espectral por subportadora é de 12,5 GHz, que é o valor adotado pelos ROADMs com tecnologia *flexible-grid* (ROSÁRIO, 2014) seguindo convenções já estabelecidas (REIXATS, 2014).

Tabela 2.1: Formatos de Modulação.

Formato de Modulação	Taxa de Símbolo	Distância Máxima (km)	Eficiência Espectral (bit/s/Hz)	Banda por Subportadora (GHz)
16QAM	1/4	500	4	12,5
8QAM	1/3	1000	3	12,5
QPSK	1/2	2000	2	12,5
BPSK	1	4000	1	12,5

Se a taxa de transferência depende da distância do caminho, então a banda alocada para uma dada demanda também irá depender do caminho pelo qual a demanda foi roteada. Suponha que cada subportadora ou SF ocupe uma banda de Ω GHz e uma demanda de v Gbit/s seja roteada com um formato de modulação que possui eficiência igual a η Gbit/s por GHz.

A Fórmula (2.4) dá o número p de SFs necessários para atender à conexão, caso o formato de modulação em questão seja usado. Já que o quociente no lado direito da equação pode ser uma fração, ela deve ser discretizada para o inteiro superior mais próximo, para que p seja um número inteiro. O operador *ceiling*, representado pelos colchetes modificados, $\lceil \cdot \rceil$, garante esta condição.

$$p = \left\lceil \frac{v}{\Omega \times \eta} \right\rceil. \quad (2.4)$$

2.2 O PROBLEMA DO RMSA

Para ilustrar um caso simples de resolução do RMSA, vamos considerar uma rede óptica hipotética, em anel, com quatro nós e quatro enlaces. Assume-se que cada nó possui transceivers, transponders e dispositivos ROADM para lidar com a transmissão/recepção necessária.

O tráfego a ser roteado nesta rede vai depender da demanda de tráfego de cada nó para cada nó, que é representada por meio de uma matriz, chamada de matriz de tráfego

$V = \{v^{ij} \in \mathbb{N}^{4 \times 4}\}$, cujo elemento v^{ij} da linha i e coluna j é a demanda de tráfego em Gbps do nó i para o nó j . Quando v^{ij} é igual a zero, significa a ausência de demanda de tráfego, como é o caso da diagonal principal, já que não existe demanda de um nó para ele mesmo, que deva ser tratada pela rede.

$$V = \begin{bmatrix} 0 & 50 & 50 & 25 \\ 70 & 0 & 65 & 30 \\ 60 & 120 & 0 & 45 \\ 100 & 35 & 150 & 0 \end{bmatrix}.$$

A Figura 2.4 mostra a resolução do RMSA para a rede em anel, usando uma heurística. Os nós e enlaces são desenhados em preto. Os nós estão numerados de 1 a 4. As setas curvadas simbolizam o caminho que uma demanda de tráfego específica deve viajar, a partir do nó de origem, para chegar ao nó de destino. Os oito enlaces e_{ij} e seus conjuntos de *slots* de frequência estão localizados à direita e à esquerda da rede anel. As caixas preenchidas de azul são os SFs que foram atribuídos para a demanda de tráfego v^{ij} , indicada logo acima da respectiva caixa. As caixas coloridas de cinza, sem indicação, são consideradas bandas de guarda.

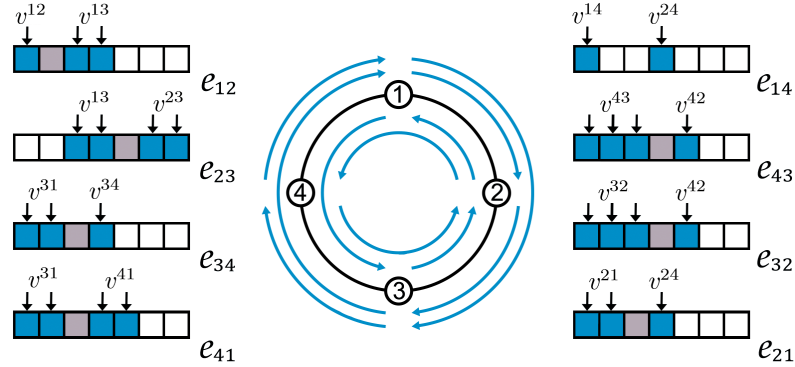


Figura 2.4: Exemplo de resolução do RMSA para uma rede anel de quatro enlaces.

Assuma que o algoritmo que resolve o RMSA para a rede anel é uma heurística que sempre escolhe o caminho mais curto para rotear a demanda que surgir. Imagine que, para este algoritmo, os formatos de modulação 8QAM e 16QAM possam ser usados e que todos os enlaces da rede têm 500 km de comprimento. Também, a política que dita a posição de alocação de cada SF dá preferência para os primeiros *slots* (*First Fit*) para alocar a carga (SFs). Considere também que o número de bandas de guarda entre conexões alocadas é igual a um.

A heurística deve ter uma regra para decidir a ordem de alocação das demandas pertencentes à matriz V . Neste exemplo, os elementos serão escolhidos da esquerda para direita e de cima para baixo (convenção de leitura). O primeiro elemento válido para alocação é a demanda de 50 Gbit/s do nó 1 para o nó 2. O caminho mais curto usa o enlace e_{12} . Já que a distância total é igual a 500 km, então é possível usar 16QAM como formato de modulação. Assim, aplicando a Fórmula (2.4), um único *slot* é alocado no enlace e_{12} .

A próxima demanda é do nó 1 para o nó 3, também de 50 Gbit/s. Existem dois caminhos mais curtos possíveis, por isso, a heurística pode escolher qualquer um dos dois caminhos. Neste caso, o caminho do nó 1 para o nó 2 para o nó 3 é escolhido. A distância total é de 1000 km, então, o melhor formato de modulação que pode ser usado é o 8QAM. Assim, dois SFs contíguos devem ser alocados de forma contínua dentro dos enlaces e_{12} e e_{23} . Note, na Figura 2.4 que a demanda v_{13} teve que ser alocada depois da banda de guarda no enlace e_{12} . Por esta razão, no enlace e_{23} , os primeiro dois *slots* não foram ocupados para obedecer a restrição da continuidade, que força o algoritmo a alocar as mesmas subportadoras para todos os enlaces, no caminho estabelecido para aquela conexão.

Se os elementos restantes da matriz de tráfego continuarem seguindo este padrão de roteamento e alocação, ao final, a resolução do RMSA produzido pela heurística como mostrado na Figura 2.4. Note que, no enlace e_{23} , apenas o primeiro *slot* pode ser ocupado, pois o segundo *slot* deve ser usado como banda de guarda. Este fenômeno, onde SFs dispersos se tornam virtualmente inutilizados ou têm uma probabilidade muito baixa de serem aproveitados é conhecido como fragmentação de espectro (AMAR D., 2015) e é uma consequência inevitável das restrições da continuidade e contiguidade.

O exemplo apresentado nesta seção teve como entrada apenas uma matriz de tráfego com a taxa de transmissão desejada, em Gbps, de cada nó para cada nó. O algoritmo recebeu a matriz de tráfego e, após a alocação de todas as demandas, o RMSA foi finalizado. Essa abordagem é tida como estática ou *Static Lightpath Establishment* (SLE) e não é uma situação realística, já que a demanda exigida de uma rede não é conhecida e atendida em um passo apenas. Em situações reais, o tráfego é considerado dinâmico e conexões novas podem surgir ou conexões já estabelecidas podem ser removidas ao longo do tempo e este tema é abordado na seção seguinte.

2.3 TRÁFEGO INCREMENTAL

Hoje, a maioria dos algoritmos de RMSA são baseados no escalonamento de múltiplas subportadoras de frequência (SHADMAND; SHIKH-BAHAEI, 2010). Em que os SFs são reservados para um usuário, por um período de tempo e depois liberados. Estes algoritmos funcionam com estabelecimento dinâmico de conexões e o RMSA é resolvido para manter o QoS em um nível aceitável, entregando uma conexão dentro de um certo limiar de latência. O estabelecimento dinâmico de tráfego é mais adequado para avaliar algoritmos de tempo real e fazer simulação de QoS.

Como a demanda de tráfego é gerada periodicamente por usuários nos terminais na rede, o tráfego agregado irá crescer continuamente e, probabilisticamente, é esperado que uma grande porção dos usuários irá gerar tráfego no mesmo momento, denominado "*rush hour*", e a rede deve ser projetada para aguentar este período de explosão de tráfego (LELAND et al., 1993; IYER; SINGH, 2018). A resolução do RMSA para tráfego incremental é idealizada para manter banda espectral livre para o futuro e manter a probabilidade de bloqueio dentro de um limiar de QoS.

Os modelos aqui propostos foram criados para fazer o estabelecimento incremental de conexões, onde novas conexões aparecem na forma de uma matriz $V = \{v^{ij} \in \mathbb{N}^{N \times N}\}$

que representa o tráfego adicional, em Gbps, que deve ser alocado (N é o número de nós da rede). A rede, então, deve resolver o RMSA para esta matriz de tráfego incremental e as conexões estabelecidas são mantidas indefinidamente.

Se considerarmos o conjunto de todas as matrizes de tráfego a serem simuladas, teremos $V^t = \{v^{ijt} \in \mathbb{N}^{N \times N \times T}\}$, $t = \{1, 2, 3, \dots, T\}$, onde N é o número de nós na rede e T é o número de matrizes de tráfego incremental simuladas. O período de alocação referente a cada V^t é denominado de período de tráfego incremental, período de tempo ou simplesmente período incremental. No t -ésimo período de tráfego incremental, a matriz V^t deve ser processada independentemente, usando os recursos disponíveis nos enlaces.

Como exemplo de resolução de RMSA para tráfego incremental usando a formulação que é apresentada no Capítulo 3 sobre a rede anel de quatro nós e duas matrizes de tráfego incremental V^1 e V^2 , obtemos o resultado ilustrado na Figura 2.5.

$$V^1 = \begin{bmatrix} 0 & 50 & 50 & 25 \\ 70 & 0 & 65 & 30 \\ 60 & 120 & 0 & 45 \\ 100 & 35 & 150 & 0 \end{bmatrix}, V^2 = \begin{bmatrix} 0 & 0 & 50 & 65 \\ 80 & 0 & 0 & 0 \\ 70 & 30 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}.$$

A matriz V^1 é igual à matriz usada no exemplo da Secção 2.2, assim como a rede simulada é a mesma. No entanto, há uma matriz de tráfego incremental V^2 , sendo $T = 2$. Para este caso, o RMSA, que é incremental, não é feito por uma heurística, assim como é imaginado no exemplo da Secção 2.2. A Figura 2.5 foi montada com os resultados de uma simulação com uma formulação de MILP de forma a reduzir o máximo SF consumido.

Neste exemplo, a primeira matriz de tráfego V^1 é alocada de forma a minimizar o o máximo SF consumido. Os SFs referentes a V^1 têm cor azul. Após a alocação da matriz V^1 a matriz V^2 é alocada, podendo consumir apenas os SFs restantes. Os *slots* laranja representam as demandas alocadas referentes a V^2 .

A disposição dos *slots* para V^1 é diferente da configuração da Figura 2.4. Isto por que a forma como o RMSA foi resolvido é diferente para o caso da Figura 2.5. O máximo SF usado por V^1 na Figura 2.5 foi o 5º, enquanto que na Figura 2.4 foi o 7º, mostrando a superioridade de otimização da formulação de MILP em minimizar o máximo SF em uso.

Os passos necessários de roteamento e alocação que geraram a configuração da Figura 2.5 não podem ser explicados, como foi feito para a Figura 2.4. Isto por que a resolução do RMSA com a formulação de MILP é feita com a minimização numérica de uma função objetivo submetida a uma série de restrições, por intermédio de um algoritmo como o IBM ILOG CPLEX. A descrição da formulação de MILP aparece no capítulo seguinte.

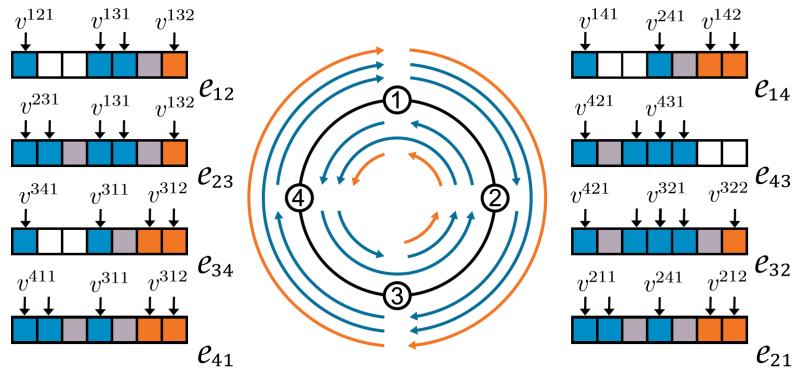


Figura 2.5: Resultado da simulação para duas matrizes de tráfego incremental, relativas a dois períodos incrementais consecutivos, $t = 1$, representado em azul e $t = 2$ representado em laranja.

MÉTODOS DE RESOLUÇÃO DO RMSA INCREMENTAL

Este capítulo apresenta as estratégias para resolução do RMSA incremental, propostas pelos autores. A Secção 3.1 mostra a formulação matemática de MILP, a Secção 3.2 mostra as heurísticas procedurais e a Secção 3.3 faz uma comparação entre os métodos e simulações de probabilidade bloqueio.

3.1 FORMULAÇÃO DE PROGRAMAÇÃO LINEAR

Algumas formulações de MILP foram propostas para resolver o RMSA, mas nenhuma foi proposta para lidar com o RMSA para tráfego incremental. Os autores deste trabalho introduziram uma formulação de MILP para gerar soluções otimizadas para T matrizes de tráfego, que são tratadas sequencialmente. Cada matriz de tráfego equivalendo a um período de incremento de tráfego.

Após definir a nomenclatura na Subsecção 3.1.1, os parâmetros de entrada na Subsecção 3.1.2 e variáveis usadas na formulação, na Subsecção 3.1.3, o modelo completo de MILP é apresentado na Subsecção 3.1.4.

3.1.1 Notações

- (i, j) e (k, u) denotam pares ordenados origem-destino, em que i e k são os nós remetentes e j e u são os nós destinatários.
- m e n denotam nós terminais de um enlace físico da rede.
- z denota o formato de modulação do conjunto M com todos os formatos de modulação disponíveis.
- t, t_1 and t_2 denotam um período de tempo do conjunto T , em que t_1 e t_2 podem ser usados para identificar dois períodos distintos de tráfego incremental.

3.1.2 Parâmetros de Entrada

- $G = (N, E)$: Um grafo com conjunto de nós N , em que cada nó é associado com um nó da rede física, e um conjunto de arestas E em que cada aresta é associada a um enlace físico da rede.
- B : Banda de guarda de filtragem, que representa o espaçamento mínimo de espectro entre as bandas de frequência.
- v^{ijt} : demanda de tráfego, em Gbit/s, do nó i para o nó j no período incremental t . É um elemento da matriz de tráfego V^t .
- Ω : Comprimento do *slot* em GHz.
- d_{mn} : Distância entre os nós m e n na topologia física.
- η_z : Eficiência espectral do formato de modulação z , em que $z \in M$.
- d_z : Alcance máximo da conexão usando formato de modulação z , em que $z \in M$.
- χ : Um número muito grande.

3.1.3 Variáveis

- p^{ijt} : Um valor inteiro que quantifica o número de SFs ocupados, por uma conexão do nó i para o nó j do período incremental t , na topologia física.
- P_{mn}^{ijt} : Número de SFs ocupados no enlace $m-n$, pela conexão do nó i para o nó j do período t .
- A_{mn}^{ijt} : Variável binária que indica se uma conexão do nó i para o nó j do período t existe e passa pelo enlace $m-n$. A_{mn}^{ijt} é igual a 1 se $P_{mn}^{ijt} > 0$; e igual a 0 se $P_{mn}^{ijt} = 0$.
- λ_z^{ijt} : Variável binária que indica se uma conexão do nó i para o nó j do período t emprega o formato de modulação z .
- S^{ijt} : Uma variável inteira que denota o índice do primeiro SF da conexão $i-j$ do período t .
- $W_{kut_2}^{ijt_1}$: Variável binária, igual a 1 se o índice do SF inicial da conexão $i-j$ do período t_1 é menor que o índice do SF inicial da conexão $k-u$ do período t_2 . (i.e. $S_{ijt_1} < S_{kut_2}$), e 0 caso contrário.
- c_t : Índice máximo dos SFs usados pelas conexões do período t .

3.1.4 Descrição Matemática da Formulação

A formulação proposta nesta dissertação é baseada em conceitos descritos na Secção 2.3, onde diferentes períodos de tráfego são usados e a solução otimizada é encontrada. A formulação garante, pela minimização da função objetivo, que os tráfegos referentes aos primeiros períodos incrementais sejam alocados antes que os períodos seguintes sejam considerados, assim como descrito na Secção 2.3.

- *Função Objetivo:*

$$\text{Minimizar : } \sum_{t \in T} \sum_{(i,j) \in D} (\theta^t S^{ijt} + p^{ijt}), \quad (3.1)$$

em que,

$$\theta^t = \sum_t^{|T|} \sum_{(i,j) \in D} \frac{v^{ijt}}{\Omega \cdot \min(\eta_z)}. \quad (3.2)$$

- *Restrições de Balanceamento de Fluxo:*

$$\sum_n P_{mn}^{ijt} - \sum_n P_{nm}^{ijt} = \begin{cases} p^{ijt} & m = i, \\ -p^{ijt} & m = j, \\ 0 & m \neq i, j, \end{cases} \quad (3.3)$$

$$A_{mn}^{ijt} \geq \frac{P_{mn}^{ijt}}{\chi}, \quad (3.4)$$

$$A_{mn}^{ijt} + A_{ml}^{ijt} \leq 1, \quad (3.5)$$

$$\forall (i, j) \in D, t \in T, (m, n) \in E, (m, l) \in E.$$

- *Restrições de Formato de Modulação:*

$$p^{ijt} \geq \left(\frac{v^{ijt}}{\Omega \cdot \eta_z} \right) - (1 - \lambda_z^{ijt}) \cdot \chi, \quad (3.6)$$

$$\forall (i, j) \in D, t \in T, z \in M.$$

$$\sum_z \lambda_z^{ijt} \leq 1, \quad (3.7)$$

$$p^{ijt} \leq \chi \cdot \sum_z \lambda_z^{ijt}, \quad (3.8)$$

$$\sum_z \lambda_z^{ijt} \geq \frac{v^{ijt}}{\chi}, \quad (3.9)$$

$$\chi \cdot v^{ijt} \geq \sum_z \lambda_z^{ijt}, \quad (3.10)$$

$$\sum_{mn} A_{mn}^{ijt} \cdot d_{mn} \leq \sum_z d_z \cdot \lambda_z^{ijt}, \quad (3.11)$$

$$\forall (i, j) \in D, t \in T, z \in M.$$

- *Restrições de Continuidade e Contiguidade de Espectro:*

$$W_{kut_2}^{ijt_1} + W_{ijt_1}^{kut_2} \leq 1, \quad (3.12)$$

$$W_{kut_2}^{ijt_1} + W_{ijt_1}^{kut_2} \geq [(A_{mn}^{ijt_1} + A_{mn}^{kut_2}) - 1], \quad (3.13)$$

$$S^{ijt_1} + p^{ijt_1} + B \leq S^{kut_2} + \chi \cdot [1 - W_{kut_2}^{ijt_1}], \quad (3.14)$$

$$S^{kut_2} + p^{kut_2} + B \leq S^{ijt_1} + \chi \cdot [1 - W_{ijt_1}^{kut_2}], \quad (3.15)$$

$$\forall (m, n) \in E, (i, j) \in D, t_1 \in T, (k, u) \in D, t_2 \in T, \\ (i, j, t_1) \neq (k, u, t_2).$$

- *Outras Restrições:*

$$\sum_{mn} A_{mn}^{ijt} \geq \frac{p_{ijt}}{\chi} \quad \forall (i, j) \in D, t \in T, \quad (3.16)$$

$$c_t \geq S^{ijt} + p^{ijt} \quad \forall (i, j) \in D, t \in T, \quad (3.17)$$

$$S^{ijt} \geq 0, \quad p^{ijt} \geq 0 \quad \forall (i, j) \in D, t \in T. \quad (3.18)$$

A Função Objetivo (3.1) é posta para minimizar a soma de todos os índices dos primeiros *slots* usados por cada demanda para todos os períodos de tempo: S^{ijt} e a soma de da carga roteada p^{ijt} para todas as demandas (i, j) do período de tempo t . A minimização de p^{ijt} é obtida sempre escolhendo o formato de modulação disponível com maior eficiência espectral. Note, no entanto, que diferentes graus de importância são colocados na minimização dos índices dos primeiros *slots* S^{ijt} para diferentes períodos incrementais. Para cada período incremental t , a variável S^{ijt} é multiplicada pela constante θ^t , dada pela Equação (3.2).

Dada a natureza incremental do problema, é necessário que a matriz de tráfego do período t_1 seja alocada antes da matriz de tráfego do período t_2 ($t_2 > t_1$). Para isso ser verdade, $\theta^{t_1} > \theta^{t_2}$. Para assegurar que isso irá acontecer sobre qualquer circunstância,

definimos θ^t como a soma agregada do número máximo de SFs requeridos pela matriz de tráfego do período t até o último período T . Isso garante que $\theta^{t_1} > \theta^{t_2}$ para $t_2 > t_1$.

Essa função objetivo foi escolhida para minimizar o máximo *slot* ocupado pela demanda da primeira matriz de tráfego sem levar em consideração os períodos incrementais seguintes. Então, após a primeira matriz de tráfego ser alocada, as demandas do segundo período incremental são alocadas, minimizando o máximo *slot* usado, dentro dos recursos restantes após a alocação da primeira matriz de tráfego e assim sucessivamente. Desta forma, o tráfego das matrizes precedentes são alocadas de forma otimizada sem considerar o tráfego das matrizes subsequentes, como aconteceria em um cenário de tráfego incremental.

Para manter o roteamento de tráfego apropriado na topologia da rede, restrições de balanceamento de fluxo são definidas, como a Restrição (3.3), que define o fluxo de saída para o todos os nós. No lado esquerdo da equação, temos a soma de todos os tráfegos saindo do nó m (P_{mn}^{ijt}) menos a soma de todos os tráfegos entrando no nó m (P_{nm}^{ijt}). Caso m seja o nó remetente da conexão $i - j$, então $m = i$ e o fluxo de saída é positivo e igual a p^{ijt} , em número de *slots*. Caso o nó m seja o destinatário, então o fluxo de saída é negativo e igual a $-p^{ijt}$, em número de *slots*. Caso o nó m não seja nem remetente nem destinatário ($m \neq i, j$), então este nó é intermediário e todo tráfego que chega deve sair, totalizando um fluxo nulo.

Restrição (3.4) define o indicador binário A_{mn}^{ijt} de P_{mn}^{ijt} . Se $P_{mn}^{ijt} = 0$, então $A_{mn}^{ijt} = 0$ e se $P_{mn}^{ijt} > 0$, $A_{mn}^{ijt} = 1$, pois o quociente é um número maior que 0 e A_{mn}^{ijt} deve ser binário. Note que a minimização da função objetivo fará com que $A_{mn}^{ijt} = 0$ quando $P_{mn}^{ijt} = 0$.

Restrição (3.5) se assegura que toda demanda (i, j) do período t só pode sair de um nó m através de um enlace. As opções de enlace de saída (m, n) e (m, l) são mutuamente exclusivos, já que sua soma deve ser menor ou igual a 1. Essa restrição garante que o tráfego não seja dividido por dois enlaces, até que chegue ao seu destino.

A Restrição (3.6) define o tráfego, em número de SFs, p^{ijt} , para cada conexão (i, j) do período t . A partir da respectiva demanda de tráfego v^{ijt} , em Gbps, e da eficiência espectral: η_z . Se o formato de modulação z for escolhido para a conexão (i, j) do período t , λ_z^{ijt} será igual a 1 e a Restrição (3.6) definirá cada p^{ijt} usando a Equação (2.4). No caso em que λ_z^{ijt} é 0 (i.e. z não foi escolhido para a tripla $i - j - t$), o resultado é que $p^{ijt} \geq -\infty$ e o valor 0 sempre será escolhido para minimizar o número de SFs, dada a função objetivo.

As restrições seguintes são postas para assegurar o gerenciamento correto dos formatos de modulação. Restrição (3.7) assegura que só existirá um formato de modulação sendo usado para uma demanda (i, j, t) . Restrição (3.8) impede p^{ijt} de assumir qualquer valor positivo quando o lado direito da inequação é igual a 0, o que significa que nenhum formato de modulação pode ser selecionado para a demanda (i, j, t) . Isso pode ocorrer quando nenhum formato de modulação da simulação não pôde suportar a distância do caminho da origem para o destino.

Para uma demanda (i, j, t) , Restrição (3.9) irá forçar pelo menos um dos λ_z^{ijt} a ser 1, quando v^{ijt} for maior que 0. Caso a demanda não exista, isso é: $v^{ijt} = 0$, esta restrição não fará efeito.

Restrição (3.10) irá garantir que se não existe demanda e $v^{ijt} = 0$, então nenhum

formato de modulação será usado, pois o lado direito da inequação será forçado a ser 0 também. Caso $v^{ijt} > 0$, esta inequação não terá efeito.

Para gerenciar as restrições de alcance dos formatos de modulação, a Restrição (3.11) é definida para que a distância total percorrida pela demanda (i, j, t) , calculada com o lado esquerdo da inequação, seja menor ou igual à distância máxima suportada pelo formato de modulação z , calculada com o lado direito da inequação.

Restrição (3.12) mantém o valor de $W_{kut_2}^{ijt_1} + W_{ijt_1}^{kut_2}$ menor ou igual a 1, já que $W_{kut_2}^{ijt_1}$ e $W_{ijt_1}^{kut_2}$ não podem ambos serem iguais a 1, pela sua definição. Restrição (3.13) garante que, entre as demandas de tráfego (i, j, t_1) e (i, j, t_2) , uma delas será alocada antes da outra, de acordo com a definição da variável $W_{kut_2}^{ijt_1}$. Se as duas demandas de tráfego atravessam um enlace em comum, o lado direito da Restrição (3.13) será igual a 1, forçando a condição $W_{kut_2}^{ijt_1} + W_{ijt_1}^{kut_2} = 1$, significando que ou (i, j, t_1) ou (i, j, t_2) vai iniciar antes do outro. Se as demandas não passam pelo mesmo enlace, o lado direito da inequação será igual a 0 e $W_{kut_2}^{ijt_1} + W_{ijt_1}^{kut_2}$ pode ser igual a 0 ou a 1, não importando se eles um se inicia antes do outro ou não.

Restrições (3.14) – (3.15) existem para assegurar que não haverá sobreposição de espectro entre todas as demandas sendo roteadas, nos enlaces da rede. Para duas demandas diferentes (i, j, t_1) e (i, j, t_2) , se elas estão no mesmo enlace, ou $W_{kut_2}^{ijt_1}$ ou $W_{ijt_1}^{kut_2}$ estará ativado. Se $W_{kut_2}^{ijt_1}$ for o ativo, Restrição (3.14) se certificará que o *slot* inicial da demanda (i, j, t_1) mais os *slots* alocados p^{ijt_1} mais a banda de guarda, B será numericamente menor que o *slot* inicial da demanda (k, u, t_2) . Neste caso, $W_{ijt_1}^{kut_2} = 0$ e a Restrição (3.15) não terá efeito. Caso, (k, u, t_2) inicie antes de (i, j, t_1) , as funções da Restrição (3.14) e da Restrição (3.15) são invertidas. Considere a situação em que nenhum enlace é compartilhado, $W_{kut_2}^{ijt_1}$ e $W_{ijt_1}^{kut_2}$ serão 0 e ambas as restrições serão triviais.

A Restrição 3.16 garante que, se houver tráfego, p^{ijt} será igual a 1, logo, devido ao lado esquerdo da inequação, para um dado (i, j, t) , pelo menos um A_{mn}^{ijt} deverá ser igual a 1. Garantido isto, as Restrições 3.4 e 3.3 se certificarão que o roteamento seja feito corretamente. A Restrição 3.17 define a variável c_t como o SF máximo usado no período t .

3.2 HEURÍSTICAS DE RMSA INCREMENTAL

Formulações de MILP são constantemente usadas para otimizar a resolução do RMSA em redes ópticas. No caso apresentado na Secção 3.1, uma formulação completa para RMSA incremental de tráfego multiperíodo foi proposta. Uma das desvantagens de formulações de MILP é que o tempo de execução cresce exponencialmente com o número de variável envolvidas no problema. Quando considerado o número de nós, enlaces, formatos de modulação e períodos incrementais, o tempo de processamento pode se tornar inviável para redes muito grandes.

Para resolver o RMSA para tráfego incremental multiperíodo em tempo mais curto, nós propomos duas heurísticas de alocação, baseadas em heurísticas clássicas de RSA (WANG; CAO; PAN, 2011) de roteamento por caminho mais curto e de roteamento por balanceamento de carga. As heurísticas propostas neste capítulo foram adaptadas para lidar com diferentes formatos de modulação e tráfego incremental.

Para as heurísticas, a rede é alimentada a um algoritmo simulador de eventos discretos, no qual o evento básico é fornecer conexões, vindas de uma matriz de tráfego V^t , para a rede. O simulador deve obter o estado da rede após a resolução do RMSA para o evento em questão e memorizar este estado para ser usado no evento seguinte. Cada evento é tido como um período de tráfego incremental, onde uma matriz de tráfego V^t deve ser atendida e os recursos da rede serão ocupados.

Para a matriz de entrada V^t , as requisições v^{ijt} são escolhidas de forma ordenada pelo algoritmo e fornecidas a uma heurística de RMSA. Se os recursos da rede não forem suficientes para satisfazer à requisição, a conexão é rejeitada e perdida para sempre. Se os recursos disponíveis no estado atual forem suficientes para atender à conexão, ela é alocada e o estado da rede é atualizado e memorizado.

o algoritmo mantém gravado a métrica de probabilidade de bloqueio, definida para um período incremental t como a porcentagem de tráfego perdido em relação a todo tráfego oferecido durante o período t .

Algumas variáveis adicionais e funções são apresentadas para descrever as heurísticas propostas a seguir.

3.2.1 Variáveis e Funções

- G : Estrutura que representa o grafo da rede, contendo todos os enlaces e_{mn} . Esta estrutura também fica encarregada de implementar o número total de SFs usados $\|e_{mn}\|$ no enlace e_{mn} .
- G_k^{ijt} : Estrutura que representa o k -ésimo caminho mais curto do nó i até o nó j . Esta estrutura contém os enlaces e_{mn} que pertencem a esta rota, a distância total para a conexão d_k^{ijt} , a carga total $\|e_{mn}\|$ em cada um dos enlaces, o formato de modulação mais eficiente para o caminho η_k^{ijt} e número de SFs consumidos pela demanda v^{ijt} usando η_k^{ijt} , isto é: p_k^{ijt} .
- G^{ijt} : Estrutura que contém k caminhos G_k^{ijt} .
- G_o : Estrutura contendo o caminho escolhido pela heurística de RMSA, junto com sua distância total d_o , carga total máxima em seus enlaces L_o , formato de modulação mais eficiente para o caminho η_o e número de SFs consumidos pela demanda v^{ijt} usando η_o , isto é: p_o . G_o pode ser imaginado como um subgrafo da rede G .
- s : Variável booleana que assume os estados **verdadeiro** ou **falso**, como resultado da função A .
- x : Valor numérico para contar as alocações bem sucedidas.
- K : Função para calcular os k caminhos mais curtos do nó i para o nó j , a partir da demanda v^{ij} e da estrutura da rede, G .
- A : Função que recebe uma demanda de tráfego v^{ij} , o conjunto de enlaces G e a estrutura com todos os formatos de modulação disponíveis M e tenta realizar o RMSA.

- F : Função de *First Fit*. Recebe a rota escolhida p^o e a carga relativa à rota escolhida p^o e aloca os SFs, usando a política de alocação conhecida como *First Fit*.

3.2.2 Algoritmo Base

Aqui é apresentado o algoritmo base para gerenciar o tráfego incremental (Algoritmo 1) e computar a probabilidade de bloqueio para cada V^t . Este algoritmo usa uma das duas heurísticas de RMSA (Algoritmo 2 ou Algoritmo 3). o objetivo das heurísticas é fazer a alocação dos recursos para as demandas de tráfego fornecidas pelo Algoritmo 1. Para conseguir este feito, o número de alocações bem sucedidas, a cada período, incremental é contado usando a variável x , inicializada como zero (Linha 1) do Algoritmo 1. Então, para cada demanda de tráfego v^{ijt} dentro da matriz de tráfego V^t (Linha 2), A função A é usada para fazer a alocação, onde a função A é uma das heurísticas de RMSA. Se a alocação for bem sucedida, a função A retornará um valor **verdadeiro** para a variável s e a variável x será incrementada por um (Linhas 3-6). Após a alocação de V^t , como saída, o Algoritmo 1 fornece o número de alocações completadas x e a rede G com um novo estado de ocupação de recursos após o RMSA ter sido completo para o período incremental t .

Dado que G é uma das saídas do algoritmo, após o processamento da matriz V^t do período t , a mesma rede G pode ser usada novamente com uma nova matriz de tráfego, assim, as novas demandas terão que ser servidas usando os SFs livres após a alocação da primeira matriz de tráfego.

Algoritmo 1 Algoritmo Base

Entrada: V^t, G, M

```

1: Inicializar:  $x = 0, \Omega = 12, 5$ 
2: for  $\forall v^{ijt} \in V^t$  do
3:    $s \leftarrow A(v^{ijt}, G, M)$ ;
4:   if  $s = \text{true}$ , then
5:      $x \leftarrow x + 1$ 
6:   end if
7: end for
```

Saída: x, G

3.2.3 Caminho Mais Curto com Modulação para Carga Mínima (SPMLM)

O Algoritmo 2 é a primeira heurística apresentada e seu modo de operação é simples. Esta heurística irá sempre escolher o caminho mais curto G_o , obtido usando o algoritmo de Dijkstra, que fornece os enlaces e_{mn} que compõe o caminho mais curto para a demanda v^{ijt} (Linha 2). A distância total d_o , do caminho escolhido, é dado como a soma de todas as distâncias de cada enlace $m - n$ compondo G_o (Linha 3). A melhor eficiência η_o é então escolhida como a maior eficiência η_z disponível, dado que a distância total d_o do caminho escolhido, é menor ou igual à distância máxima aceitável d_z para aquele formato de modulação (Linha 4).

Depois disso, o tráfego, medido em número de *slots*, é computado usando a demanda de transmissão, em Gbit/s, a eficiência do caminho escolhido e a banda por SF Ω somada à banda de guarda B (Linha 5). No final, a política de alocação *First Fit*, representada pela função F (Linha 6). Se assume que a função F retorna **verdadeiro** para s quando a alocação é bem sucedida e **falso** caso contrário. O Algoritmo 2 foi nomeado de Caminho Mais Curto com Modulação para Carga Mínima ou *Shortest Path with Minimum Load Modulation* (SPMLM).

Algoritmo 2 Heurística (SPMLM)

Entrada: v^{ijt} , G , M

- 1: Inicializar: G_o, s
- 2: $G_o \leftarrow \text{Dijkstra}(v^{ijt}, G)$
- 3: $d_o \leftarrow \sum d_{mn} \in G_o$
- 4: $\eta_o \leftarrow \max(\eta_z \in M : d_z \geq d_o)$
- 5: $p_o \leftarrow \left\lceil \frac{v^{ijt}}{\Omega \times \eta_o} \right\rceil + B$
- 6: $s \leftarrow F(p_o, G_o)$

Saída: s

3.2.4 Balanceamento e Modulação para Carga Mínima (BMLM)

O Algoritmo 3 é outra heurística que pode ser usado como função A , cujo objetivo é balancear a carga em todos os enlaces da rede. De início, a estrutura com k caminhos mais curtos $\mathbf{G}^{ijt}$ e a estrutura com o caminho mais curto G_o são inicializadas. Essas estruturas serão usadas mais à frente. (Linha 1).

O segundo passo para esta heurística é calcular os k caminhos para a demanda v^{ijt} , usando a função K (Linha 2). Esta etapa pode ser feita usando um dos algoritmos de encontrar os k caminhos mais curtos, tal como o proposto em (YEN, 1971). Esses caminhos ficarão dentro da estrutura $\mathbf{G}^{ijt}$. Depois disso, para cada um dos caminhos G_k^{ijt} , dentro de $\mathbf{G}^{ijt}$ (Linha 3), a distância total d_k^{ijt} é computada para este caminho, como a soma de todas as distâncias individuais d_{mn} dos enlaces que compõem o caminho (Linha 4). Também é calculada a carga total máxima L_k^{ijt} do caminho k como valor máximo de número de *slots* ocupados dentro todos os enlaces do caminho G_k^{ijt} (Linha 5). Em seguida, a melhor eficiência η_k^{ijt} é definida dentre todos os formatos de modulação, tal que, a distância total d_k^{ijt} do caminho é menor ou igual à máxima distância aceitável para aquele formato de modulação d_z (Linha 6). Com a melhor eficiência para k -ésimo caminho definida, a carga p_k^{ijt} consumida pela demanda v^{ijt} é computada caso o caminho G_k^{ijt} seja usado, mais à banda de guarda B (Linha 7).

Com d_k^{ijt} , L_k^{ijt} , η_k^{ijt} e p_k^{ijt} adquiridos para cada G_k^{ijt} , a carga máxima do caminho escolhido L_o é calculada como a carga do caminho no qual, a soma da carga máxima L_k^{ijt} mais a respectiva demanda, em número de *slots*, $p_k^{ijt} + B$ é a menor de todas (Linha 9). Depois da menor carga total L_o ser escolhida, a menor distância total d_o é escolhida dentre todas as distâncias d_k^{ijt} dos caminhos que possuem carga máxima L_k^{ijt} igual a L_o (Linha 10). Dessa forma, o caminho mais curto com a menor carga máxima é escolhido

como G_o (Linha 11). O tráfego em número de *slots* p_o será o tráfego do caminho G_k^{ijt} , isto é, p_k^{ijt} (Linha 12). Então, a política de alocação *First Fit* é chamada para alocar a carga no caminho escolhido (Linha 6). O Algoritmo 3 foi nomeado de Balanceamento e Modulação para Carga Mínima ou *Balanced Minimum Load Modulation* (BMLM).

Algoritmo 3 Heurística (BMLM)

Entrada: v^{ijt}, G, M

- 1: Inicializar: $G_o, \mathbf{G}^{ijt}, s$
- 2: $\mathbf{G}^{ijt} \leftarrow K(v^{ijt}, G)$;
- 3: **for** $\forall G_k^{ijt} \in \mathbf{G}^{ijt}$ **do**
- 4: $d_k^{ijt} \leftarrow \sum_{e_{mn} \in G_k^{ijt}} d_{mn}$
- 5: $L_k^{ijt} \leftarrow \max(\|e_{mn}\| : e_{mn} \in G_k^{ijt})$
- 6: $\eta_k^{ijt} \leftarrow \max(\eta_z \in M : d_z \geq d_k^{ijt})$
- 7: $p_k^{ijt} \leftarrow \left\lceil \frac{v^{ijt}}{\Omega \times \eta_k^{ijt}} \right\rceil + B$
- 8: **end for**
- 9: $L_o \leftarrow L_k^{ijt} = \min(L_k^{ijt} + p_k^{ijt})$
- 10: $d_o \leftarrow \min(d_k^{ijt} : L_k^{ijt} = L_o)$
- 11: $G_o \leftarrow G_k^{ijt} : d_k^{ijt} = d_o, L_k^{ijt} = L_o$
- 12: $p_o \leftarrow p_k^{ijt}$
- 13: $s \leftarrow F(p_o, G_o)$

Saída: s

Tanto o SPMLM quanto o BMLM são heurísticas para tentar minimizar o tráfego na rede, a partir de um algoritmo predeterminado. O SPMLM escolhe sempre o caminho mais curto como forma de prover a conexão entre dois nós. O BMLM procura balancear a carga ao tentar escolher o caminho com menor carga máxima em seus enlaces. Essas são estratégias clássica usadas na literatura, mas neste trabalho elas são adaptadas para trabalhar com tráfego incremental, de forma que, ao inserir a estrutura G da rede e fornecer uma matriz de tráfego V^t , representando um período incremental de tráfego, e alimentar esses parâmetros no Algoritmo 1, ele usará uma das heurísticas que tentará fazer o RMSA. Após a alocação, a estrutura G será recebida na saída, atualizada com os SFs ocupados pelas conexões bem sucedidas. Obviamente, as heurísticas não fornecem o melhor resultado possível, como faz a formulação apresentada na Secção 3.1, mas eles são a melhor opção quando o tempo para processamento é escasso.

Na Secção 3.3 é mostrada uma pequena demonstração da superioridade da formulação de MILP em comparação às heurísticas, para redes pequenas. No entanto, quando as simulações são feitas com redes grandes, as heurísticas são usadas exclusivamente, devido a limitações do *hardware*.

3.3 RESULTADOS DO RMSA INCREMENTAL

Apresentados os métodos de resolução do RMSA incremental, esta secção faz uma comparação do consumo de capacidade consumida da formulação de MILP e das heurísticas

para redes pequenas na Subsecção 3.3.1 e o caso das redes grandes é visto na Subsecção 3.3.2 e simulações mais realistas são vistas na Subsecção 3.3.3, em termo de probabilidade de bloqueio.

3.3.1 Redes Pequenas

Para testar a funcionalidade da formulação de MILP apresentada na Secção 3.1, simulações foram realizadas para as redes com 6 nós, ilustradas na Figura 3.1. Três períodos de tráfego incremental foram simulados como 3 matrizes de tráfego V^t , $t = \{1, 2, 3\}$. Cada nó exigindo 100 Gbit/s de transmissão para cada outro nó, isto é $v^{ijt} = 100$ Gbit/s $\forall(i, j, t)$. A banda por *slot*, Ω , foi configurada como 12,5 GHz e a banda de guarda de filtro, entre as conexões foi configurada como um *slot*. É assumido que quatro formatos de modulação estarão disponíveis ($|M| = 4$). A eficiência espectral de cada formato de modulação, η_z , foi dado como $\eta_1 = 1$, $\eta_2 = 2$, $\eta_3 = 3$ e $\eta_4 = 4$ bit/s/Hz. O alcance máximo de cada conexão sob cada formato de modulação z é $d_1 = 8$ saltos, $d_2 = 4$ saltos, $d_3 = 2$ saltos e $d_4 = 1$ salto. As simulações foram feitas usando tanto a formulação de MILP apresentada na Secção 3.1 como as heurísticas apresentada na Secção 3.2.

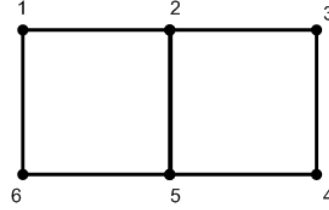


Figura 3.1: Rede pequena com seis nós, usada para as simulações nesta secção

A simulação de MILP usou o pacote de otimização IBM ILOG CPLEX (ILOG, Inc, 2019) como *solver*. A simulação foi executada em um Intel i7 3.6 GHz with 32 GByte de memória RAM, durante um período de 5 dias (120h). Os resultados são apresentados na Tabela 3.1. Lembramos que as heurísticas têm acesso ao índices dos SFs para fazer a alocação, pela função F , e os valores de c_t podem ser facilmente acessados na estrutura G .

Tabela 3.1: Simulações com MILP

Período Incremental (t)	1	2	3
	Índice do <i>slot</i> Máximo (c_t)		
MILP	20	43	66
SPMLM	28	54	78
BMLM	31	99	146

A formulação teve uma performance melhor e foi eficaz em minimizar a função objetivo e consequentemente, o máximo SF usado c_t , nos enlaces, para todos os períodos incrementais. A heurística SPMLM ficou em segundo lugar e a BMLM em terceiro. Este resultado é indicativo da ineficiência do BMLM em redes pequenas, pois o BMLM

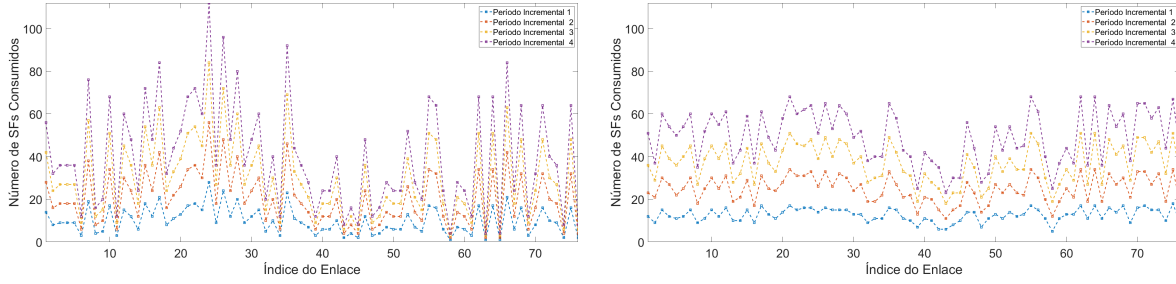


Figura 3.2: Número de SFs consumidos (Carga) por enlace da rede óptica Européia por quatro períodos de tráfego incremental. (a) SPMLM e (b) BMLM.

atinge menores índices máximos de SFs em simulações com redes maiores como é visto na Subsecção 3.3.2. Por conta do tempo de processamento, fazer mais comparações com a formulação de MILP se torna inviável para esta pesquisa, mas nas subsecções seguintes, comparações entre as heurísticas são estudadas mais profundamente.

3.3.2 Redes Grandes

Por causa da complexidade do problema, para redes grandes (com mais de 7 nós), a estratégia apresentada na Secção 3.1 pode consumir muito tempo, devido ao crescimento exponencial do tempo de processamento. Como exemplo, simulação com a rede 4 nós e 4 enlaces da Secção 2.3 levou 10 minutos de processamento. A simulação com a rede de 6 nós e 7 enlaces no exemplo da Subsecção 3.3.1 levou 5 dias. Se mais 1 nó e 2 enlaces fossem adicionados à rede de 6 nós da Figura 3.1, o tempo de processamento necessário seria de 11 dias, como testado pelos autores, mostrando a necessidade de modelos heurísticos para lidar com redes moderadamente maiores como a que é simulada nesta subsecção.

Como uma forma de mostrar as consequências de se usar as heurísticas SPMLM e o BMLM como algoritmo de RMSA ao longo do tempo, foi feita uma simulação com a rede óptica europeia (*European Optical Network*), que possui 19 nós e 76 enlaces. Foram simulados 4 períodos de tráfego incremental usando apenas o BPSK como formato de modulação. A demanda de tráfego foi de 12 Gbps para todo par de nó, nos 4 períodos incrementais, ou seja, $v^{ijt} = 12 \text{ Gbps}$ e $p^{ijt} = 1 \text{ SF}$ para todo (i, j, t) .

A Figura 3.2 mostra o número de SFs consumidos (carga) para cada m dos enlaces da rede para quatro períodos de tráfego incremental, onde (a) é a disposição da carga para o SPMLM e (b) para o BMLM. A especificação dos enlaces é feita pelo índice, cuja ordem não importa.

É notável que a distribuição de carga para o BMLM é mais balanceada relativa à do SPMLM, onde alguns enlaces têm carga excessiva e outros enlaces possuem ficam subutilizados. Isso pode levar à suposição de que o consumo de recursos com o BMLM é menor, no entanto, a Tabela 3.2 dá indícios que o BMLM pode não ser tão interessante a longo prazo.

A Tabela 3.2 mostra o índice do *slot* máximo c_t para e a carga média $\bar{L}_t$ por enlace, para ambas heurísticas de RMSA incremental. Quatro períodos de tráfego incremental aparecem.

Tabela 3.2: Simulações com Heurísticas

Heurística	Métrica	Período Incremental			
		1	2	3	4
SPMLM	Índice do <i>slot</i> Máximo (C_t)	55	111	167	223
	Carga Média (L_t)	17	35	53	71
BMLM	Índice do <i>slot</i> Máximo (C_t)	32	66	99	134
	Carga Média (L_t)	22	45	67	90

Para esta simulação, o BMLM teve menor c_t , mas a carga média $\bar{L}_t$ é maior. Isso ocorre por que o BMLM percorre caminhos mais longos como forma de buscar rotas alternativas para distribuir a carga. Como consequência, a quantidade de enlaces com *slots* ocupados será maior, aumento a carga média.

É possível fazer suposições sobre os resultados encontrados na até agora, no entanto, as simulações feitas até agora não são fiéis às condições a que uma rede óptica de *backbone* seria submetida em situações reais. As simulações da Subsecção 3.3.1 trata apenas de uma rede de seis nós, com apenas dois períodos de tráfego incremental determinístico. O exemplo desta subsecção usa uma rede grande e quatro períodos de tráfego incremental, mas as demandas de tráfego são determinísticas e crescem lentamente, onde cada demanda exige apenas um SF de demanda.

3.3.3 SPMLM vs BMLM

Como forma de fazer simulações mais realistas, para redes grandes e comparar a performance do SPMLM e do BMLM novas simulações foram feitas, desta vez usando valores variáveis de demandas de tráfego e um maior número de redes ópticas.

As heurísticas da Secção 3.2 foram aplicadas a 23 redes de *backbone* conhecidas (ROUTRAY et al., 2013) com o número de nós variando de 9 a 26. O tempo de simulação, das heurísticas, foi inferior a 5 segundos em todos os casos. Devido ao alto número de redes consideradas, a amplitude dos valores de capacidade é muito grande para representar de forma intuitiva ao longo do tempo, por isso, apresentamos as simulações periódicas de probabilidade de bloqueio, que é uma métrica normalizada para todas as redes.

Um total de 25 períodos de tráfego incremental foram simulados, alimentando 25 matrizes de tráfego $V^1, V^2, \dots, V^{25}$ para o Algoritmo 1 como V^t . As matrizes de tráfego foram geradas como matrizes de mesma dimensões que a matriz de custo da rede (N^2), em que cada elemento v^{ijt} poderia assumir um valor entre 0, 100, 200, 300 e 400 Gbit/s, com uma distribuição uniforme de probabilidade.

Para todas as redes simuladas, as matrizes de tráfego geradas foram inseridas sequencialmente como entrada para simular os períodos de tráfego incremental. O número de conexões bem sucedidas x . foi usado, relativo ao total de conexões em V^t , para calcular a probabilidade de bloqueio para cada período incremental. A probabilidade de bloqueio pode ser usada como métrica de planejamento de rede (BARRY; HUMBLET, 1995), estimativa de QoS (YOO; QIAO; DIXIT, 2000), além de fornecer informação sobre fragmentação e o estado dos recursos da rede.

Ao gerar a matriz V^t dessa forma, o tráfego total fornecido por período aumenta com o número de nós da rede, já que o número de conexões é dado por N^2 . Por isso, a capacidade dos enlaces, em numero de SFs foi variado para redes de tamanhos diferentes, evitando que as redes tenham bloqueio alto logo nos primeiros períodos de tempo. O número máximo de SFs para cada rede e a lista de redes simuladas pode ser encontrado na Tabela 3.3. Para a mesma rede, a capacidade dos enlaces foi tida constante para todas as smilações.

Para as simulações, ao invés de se usar a distância entre os enlaces, foi usado o número de saltos para encontrar os caminhos mais curtos e calcular as distâncias aceitadas pelos formatos de modulação. Os formatos de modulação usados foram os presentes na Tabela 2.1 considerando um salto igual a 500 km.

Tabela 3.3: Parâmetros de redes e simulações. A referências com os parâmetros reais das redes: <http://www.av.it.pt/anp/on/refnet2.html>

Número de Nós por Rede (Parte I)						
9	10	11	12	14	15	17
Via	BREN	Abilene	CESNET	Italiana	ACONET	Germany
-	LEARN	CompusServe	VBNS	NFSNET	-	Spain
-	RNP	-	-	-	-	-
Capacidade Máxima por Enlace						
512	512	512	512	512	512	1024
Número de Nós por Rede (Parte II)						
19	20	21	23	24	25	26
CANARIE	ARPANET	PIONIER	Bulgaria	COX	SANET	NewNet
EON	Sweden	-	-	-	-	-
Memorex	-	-	-	-	-	-
Capacidade Máxima por Enlace						
1024	1024	2048	2048	2048	2048	2048

Primeiro, as simulações foram feitas usando o SPMLM e o BMLM com apenas a quarta opção de formato de modulação da Tabela 2.1, isto é, $M = \{d_4 = 4000, \eta_4 = 1\}$. Para ilustrar os resultados, o diagrama de caixa foi escolhido, por poder mostrar a tendência geral da probabilidade de bloqueio das redes sem ser severamente afetada pelos *outliers*. Os resultados são apresentados na Figura 3.3.

A Figura 3.3 mostra dois diagramas de caixa, lado a lado. Um diagrama mostra a extensão de valores para o SPMLM e o outro, para os resultados com o BMLM. Este diagrama ilustra, para todas as redes da Tabela 3.3. As tendências para as duas heurísticas são similares, mas é possível perceber, olhando a Figura 3.3, que, a partir do terceiro período de tráfego incremental, o BMLM apresenta uma probabilidade de bloqueio maior. Apesar do BMLM ter maior probabilidade de bloqueio a partir do terceiro período, durante os dois primeiros períodos, o BMLM tem performance melhor que a do SPMLM.

O BMLM tem melhor performance inicial, com apenas um formato de modulação por que o SPMLM sempre escolhe o caminho mais curto. Isto causa uma situação em

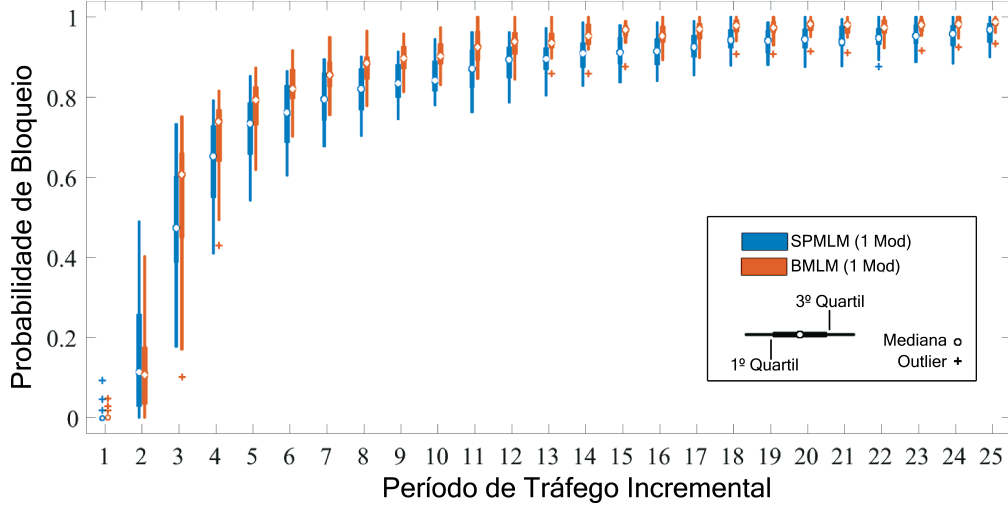


Figura 3.3: Diagrama de caixa da probabilidade de bloqueio relativo a cada período de tráfego incremental simulado. Apenas um formato de modulação foi usado, tanto para o SPMLM quanto para o BMLM.

que poucos enlaces são usados para a maioria das conexões e se tornam sobrecarregados, como consequência, fazendo com que alguns poucos enlaces comecem a bloquear tráfego mais cedo. O BMLM tenta evitar tais situações fazendo balanceamento de carga. Esta estratégia parece funcionar bem para os primeiros períodos de tempo, mas ao tentar balancear a carga, no longo prazo, o BMLM usará rotas maiores, consumindo SFs em um número maior de enlaces e aumento a carga média da rede, como visto na Tabela 3.2. Além disso, por trocar constantemente de rotas, como forma de balancear a carga, o BMLM irá criar maior fragmentação espectral. Assim, o BMLM, inevitavelmente, terá pior performance temporal, relativo ao SPMLM.

As simulações também foram feitas com todos os formatos de modulação disponíveis da Tabela 2.1: $M = \{d_1 = 500, \eta_1 = 4, d_2 = 1000, \eta_2 = 3, d_3 = 2000, \eta_3 = 2, d_4 = 4000, \eta_4 = 1\}$. Os resultados são apresentados na Figura 3.4.

Como esperado, quando comparada a performance das simulações com apenas um formato de modulação com a simulação com quatro opções de formatos de modulação, a probabilidade de bloqueio é menor para o último caso, mas há também maior disparidade para os mesmos períodos de tempo, significando que algumas arquiteturas de rede se beneficiaram mais que outras, por ter um conjunto maior de modulações disponíveis.

É notável que a performance do BMLM foi efetivamente pior que a do SPMLM para períodos posteriores, durante a simulação. Como a eficiência espectral dos caminhos mais curtos também é maior, existe um incentivo positivo em rotear as conexões por caminhos mais curtos e o BMLM sofre punições quando escolhe caminhos mais longos para balancear a carga.

Uma visão mais intuitiva do comportamento incremental da probabilidade de bloqueio para o SPMLM e para o BMLM, para duas conjecturas de modulação, é apresentada na Figura 3.5, que mostra as curvas da probabilidade de bloqueio média para as simulações anteriores.

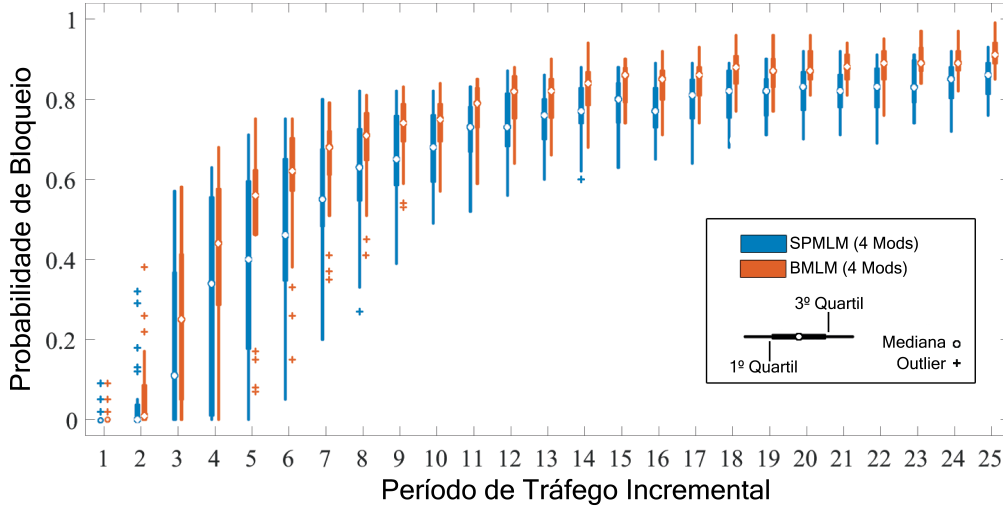


Figura 3.4: Diagrama de caixa da probabilidade de bloqueio relativo a cada período de tráfego incremental simulado. Quatro formatos de modulação foram simulados para o SPMLM e o BMLM.

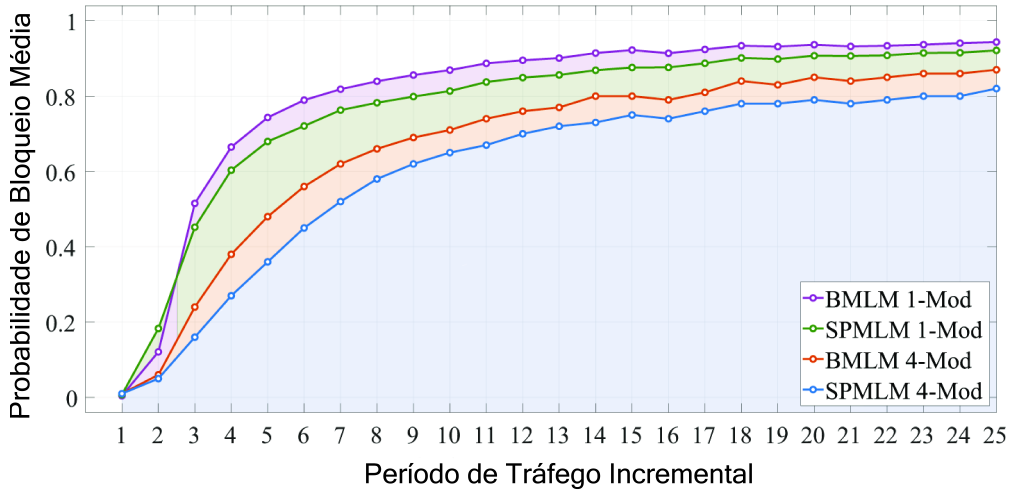


Figura 3.5: Cuvas da probabilidade de bloqueio média para os 25 períodos incrementais, e diferentes conjecturas de formatos de modulação.

Apesar dos fatores mencionados e do BMLM ter pior performance, em geral, o BMLM teve menor probabilidade de bloqueio média para algumas, como ACONET e PIONIER, provando que a heurística de RMSA pode ter resultados imprevisíveis de acordo com a topologia física da rede.

As simulações feitas aqui são mais realistas que as feitas nas subsecções anteriores, mas podem ser melhoradas usando dados de tráfego reais ao invés de tráfego gerado estocasticamente. Caso resultados similares sejam obtidos, um maior grau de confiabilidade pode ser aceito para as hipóteses feitas neste capítulo.

No Capítulo seguinte as simulações serão feitas com tráfego real de redes ópticas de *backbone*, obtidas com ajuda de modelos novos de redes neurais artificiais.

PREVISÃO DE TRÁFEGO

Recentemente, as redes neurais artificiais têm se mostrado como uma ótima ferramenta para se fazer regressão. As redes neurais recorrentes, em especial, se tornaram computacionalmente viável recentemente e foram aplicada para fazer previsão de tráfego de redes ópticas de *backbone* (ZHAO et al., 2018a; AZZOUNI; PUJOLLE, 2017a). A previsão de tráfego para redes de computadores precisa ser feita considerando as correlações do tráfego nos vários enlaces, por isso é mais preciso fazer previsões de matrizes de tráfego inteiras ao longo do tempo.

Para complementar as simulações feitas na Secção 3.3, que foi feita usando matrizes de tráfego geradas estocasticamente, os autores julgaram importante fazer simulações com dados reais de tráfego provindo de bancos de dados publicamente disponíveis. Considerando que as redes neurais artificiais podem fazer ser treinadas para fazer a previsão mencionada, elas foram empregadas de forma a confirmar a capacidade destas redes de fazer regressão com matrizes de tráfego e, como consequência, prever os dados reais para serem usados nas simulações.

A teoria de redes neurais é vista na Secção 4.1. A Secção 4.2 detalhará o procedimento de previsão de matrizes de tráfego, a Secção 4.3 explica o treinamento e a previsão. Finalmente, a Secção 4.4 mostra os resultados das simulações com os dados gerados pelas redes neurais.

4.1 REDES NEURAIIS ARTIFICIAIS

Esta secção traz uma introdução da teoria de redes neurais, começando pelo modelo mais elementar na Subsecção 4.1.1 e seguindo para as redes neurais recorrentes na Subsecção 4.1.2 e abordando memória longa de curto prazo na Subsecção 4.1.3.

4.1.1 Redes Neurais Feed-Foward

Redes neurais artificiais (RNA) são um subcampo do aprendizado de máquina, que envolve as técnicas computacionais para análise de dados em que o modelo é criado por meio de um processo iterativo (BISHOP, 2006).

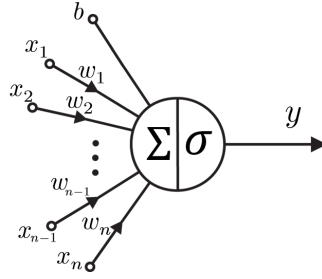


Figura 4.1: Ilustração de um neurônio artificial.

Em RNAs, o modelo é construído baseado em uma estrutura de unidades de processamento interconectadas, chamadas de neurônios. Esta estrutura é inspirada na biologia do sistema nervoso animal. Figura 4.1 mostra o modelo matemático do neurônio.

O neurônio recebe um conjunto de n entradas $x_i \mid i \in \{1, 2, 3, \dots, n-1, n\}$, chamadas de atributos ou características. Cada atributo x_i é multiplicado por um único peso w_i , então, tudo é somado com uma constante de valor b (*bias*). Depois do ponto de soma, uma função de ativação σ é aplicada na soma, fornecendo uma saída que pode ser uma classificação ou uma previsão, de acordo com a Equação (4.1).

$$y = \sigma \left(\sum_{i=1}^n w_i x_i + b \right). \quad (4.1)$$

Um único neurônio não é apropriado para lidar com a maioria das aplicações, pois ele só é capaz de fazer separação linear, quando se trata de problemas de classificação (MINSKY; PAPER, 2017). A melhor forma de encarar problemas mais complexos é fazer arranjos de muito neurônios em uma estrutura conectada composta de várias camadas, na qual a primeira camada recebe as entradas e geram saídas que serão as entradas das camadas seguintes. A Figura 4.2 mostra uma rede neural *Feed-Foward* (FFNN) completamente conectada com n entradas e duas camadas ocultas. As saídas da primeira camada oculta são as entradas da segunda camada oculta. As saídas da segunda camada oculta se tornam entradas para a camada de saída, contendo m neurônios. Neste caso em específico, a RNA usa n atributos de entrada para ter m classes de saídas. Mais de uma classe é útil se desejarmos fazer mais do que classificação binária.

Quando os neurônios estão conectados como na Figura 4.2, a saída da primeira camada é dada pela Equação (4.1), mas a saída das camadas seguintes são funções das entradas precedentes. Se indexarmos as m camadas da rede da Figura 4.2 com as variáveis l e os nós das camadas com as variáveis i, j, k e m para camada de entrada, primeira camada oculta, segunda camada oculta e camada de saída respectivamente. Então, a expressão que descreve a saída $y_k^{(2)}$ de cada neurônio, para a segunda camada oculta da rede é dada pela Equação (4.2).

$$y_k^{(2)} = \sigma \left(\sum_j w_{kj}^{(2)} \left[\sigma \left(\sum_i w_{ji}^{(1)} x_i + b_j^{(1)} \right) \right] + b_k^{(2)} \right). \quad (4.2)$$

Um procedimento similar pode ser adotado para se obter as saídas da l -ésima camada.

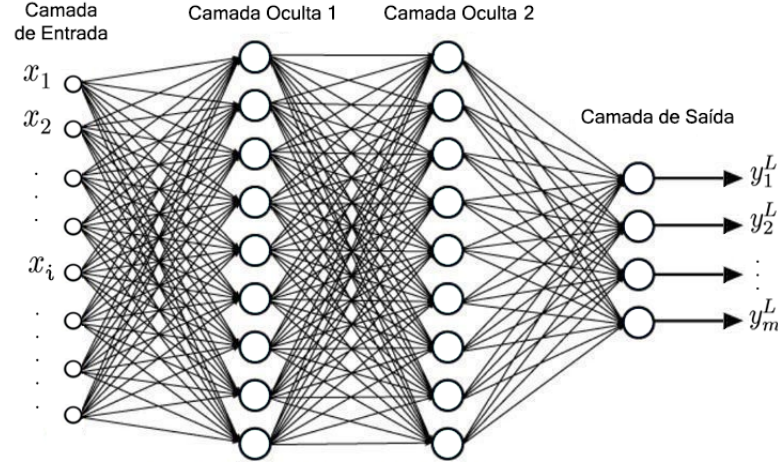


Figura 4.2: Rede Neural Artificial Feed-Forward

É fácil de perceber que, à medida que mais camadas ocultas são adicionadas à ANN, a equação que descreve a saída se torna mais complexa. Esse aspecto permite que a saída descreva dados ainda mais distintos. A equação geral para o n -ésimo neurônio da camada final L de qualquer rede ANN pode ser descrita como na Equação (4.3).

$$y_m^L = \left[\sigma \left(\sum_m w_{nm}^L \dots \left[\sigma \left(\sum_j w_{nj}^2 \left[\sigma \left(\sum_i w_{ji}^1 x_i + b_j^1 \right) \right] \right. \right. \right. \right. \\ \left. \left. \left. + b_n^2 \right) \right] \dots + b_n^L \right] = \sigma \left(\sum_n w_{nm}^L y_m^{L-1} + b_n^L \right) \quad (4.3)$$

A ANN é, então, treinada usando um processo iterativo no qual os pesos e os *bias* iniciais são escolhidos como valores aleatórios. Então, pares de entrada-saída (i.e. exemplos de treinamento) são usados para gerar a saída y , do modelo. Essa saída é comparada com o valor desejado de saída $\hat{y}$, gerando uma medida do erro da previsão. Os pesos dos neurônios são, então, ajustados para minimizar este erro. Existem diversos algoritmos propostos para o treinamento de ANNs. Dentro deles está o algoritmo tradicional de retro-propagação do erro, o algoritmo de Levenberg-Marquardt e o algoritmo de retro-propagação resiliente (RPROP).

4.1.2 Redes Neurais Recorrentes

Redes Neurais Recorrentes (RNN) são arquiteturas ANN especializadas que lidam melhor com dados com interdependência temporal (BENGIO; BOULANGER-LEWANDOWSKI; PASCANU, 2013). RNNs apresentam malhas de realimentação que são usadas para modelar dependências no tempo.

A RNN pode ser vista como um arranjo de ANNs sequenciais, cada uma atribuída para um passo de tempo específico. $t \in \{1, 2, 3, \dots, T\}$, em que T é o número de exemplos de entrada anteriores, necessários para descrever o problema de forma completa. O valor de T depende na natureza do problema a ser resolvido.

A Figura 4.3 apresenta uma arquitetura básica da RNN. As entradas definidas demarcadas por tempo $x_1, x_2, \dots, x_t$ são vetores de características que são alimentos para sua FFNN demarcada por tempo. As FFNNs usam um conjunto de parâmetro compartilhados w_{aa} , w_{ax} e w_{ya} para calcular sua saída e também para calcular um novo parâmetro a_t que é passado para o próximo passo de tempo. Dessa forma, a FFNN, a cada passo de tempo, possui informação sobre os passos de tempo anteriores, já que recebe não só x_t , mas também a_{t-1} para calcular sua saída y_t , de acordo com a Equação (4.4) e a Equação (4.5).

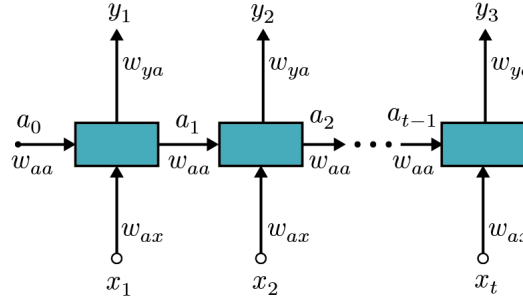


Figura 4.3: Rede Neural Recorrente como uma junção sequencial de FFNNs convencionais

$$a_t = \sigma(w_{aa}a_{t-1} + w_{ax}x_t + b_a), \quad (4.4)$$

$$y_t = \sigma(w_{ya}a_t + b_y). \quad (4.5)$$

4.1.3 Long Short-Term Memory

As RNNs apresentam ótima performance para dados temporais em sequência, mas se um grande número de passos de tempo for modelado, o algoritmo de retro-propagação ao longo do tempo não é capaz de propagar as atualizações dos pesos de volta aos primeiros passos de tempo da RNN, devido ao problema de dissipação do gradiente (HOCHREITER, 1998).

A solução para esta limitação é encontrada usando um diferente tipo de neurônio, a memória longa de curto prazo ou *Long Short-term Memory* (LSTM). (SAK; SENIOR; BEAUFAYS, 2014), designada para guardar informação de passos de tempos anteriores e levar esta informação para os últimos neurônios, através de caminhos especiais chamados de *gates*, governados pelas equações a seguir:

$$u_t = \sigma(w_{ua}a_{t-1} + w_{ux}x_t + w_{uc}c_{t-1} + b_u), \quad (4.6)$$

$$f_t = \sigma(w_{fa}a_{t-1} + w_{fx}x_t + w_{fc}c_{t-1} + b_f), \quad (4.7)$$

$$c_t = f_t \odot c_{t-1} + u_t \odot \tanh(w_{ca}a_{t-1} + w_{cx}x_t + b_c), \quad (4.8)$$

$$o_t = \sigma(w_{oa}a_{t-1} + w_{ox}x_t + w_{oc}c_{t-1} + b_o), \quad (4.9)$$

$$a_t = o_t \odot \tanh(c_t). \quad (4.10)$$

O funcionamento interno do LSTM pode ser aprendido de forma mais aprofundada através de outras fontes. Os autores recomendam os textos de Hochreiter (SAK; SENIOR; BEAUFAYS, 2014), mas o princípio básico é que a célula de memória c_{t-1} , do nó LSTM anterior pode ser passada pelos nós adjacentes, usando o *gate* de saída o_t . Este efeito pode ser resumido pelas Equações (4.8), (4.9) e (4.10). A célula de memória também pode ser esquecida ou atualizada com ajuda de f_t e u_t , respectivamente. Esses *gates* são descritos pelas Equação (4.7) e pela Equação (4.9). O diagrama na Figura 4.4 é uma ilustração simplificada do modelo de um neurônio LSTM.

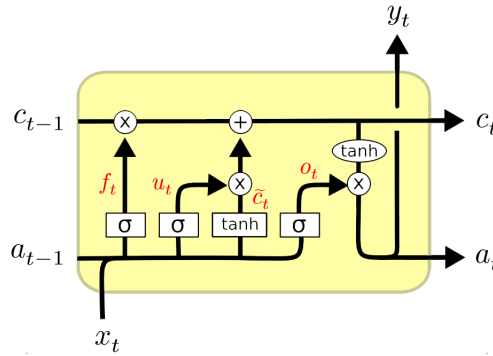


Figura 4.4: Diagrama da memória longa de curto prazo (LSTM)

4.2 PREVISÃO DE MATRIZES DE TRÁFEGO

Nesta secção, é abordado todo o procedimento de previsão de matrizes de tráfego, começando pelo banco de dados usado na Subsecção 4.2.2, apresentando em seguida os modelos de RNN propostos para a fazer a previsão de tráfego para as duas redes do banco de dados na Subsecção 4.2.2.

4.2.1 Banco de Dados

O banco dados usado para o treinamento e previsão, nesta pesquisa, foi a "European traffic matrices", que representam dados de matrizes de tráfego anonimizadas para, tanto a rede Abilene em 2004 e a rede GEANT em 2005 (BALON et al., 2007). Os dados da Abilene, com 12 nós, contêm seis meses de gravações do fluxo de tráfego instantâneo, amostrado a uma taxa de 12 amostras/minuto, totalizando 53.568 matrizes de exemplo. Os dados da rede GEANT, com 23 nós, contêm quatro meses de gravações do fluxo de tráfego instantâneo amostrado a uma taxa de 4 amostras/minuto, totalizando 11,904 matrizes de exemplo.

4.2.2 Modelos de RNNs com LSTM

Um dos problemas relacionados ao planejamento apropriado e uso efetivo de recursos, no contexto de redes ópticas, é a estimação das matrizes de tráfego (i.e. fluxo de entrada e saída de tráfego para todos os nós da rede). Ser capaz de derivar a matriz de tráfego futura de uma rede, baseando-se nas características de tráfego passado, é um desafio chave para

provedores de internet, que enfrentam problemas de congestão e sobredimensionamento da capacidade das fibras.

Este tópico foi abordado em (FELDMAN, 1981) usando média móvel integrada auto-regressiva para prever a matriz de tráfego em um futuro próximo. Em (LIU et al., 2014) foi usado previsão de matrizes de tráfego inteiras com previsão independente para cada nó individual. Este caso demonstrou resultados não satisfatórios, pois assume que não há correlação entre o fluxo de tráfego de diferentes enlaces na rede, que não é uma presunção verdadeira (WEN; ZHU, 2006).

Esta dissecação propõe modelos de RNN com LSTM, inspirados por trabalhos anteriores (AZZOUNI; PUJOLLE, 2017b; ZHAO et al., 2018b), mas com diferentes parâmetros de construção. As características de entradas da RNN são os elementos da matriz de tráfego $X = \{x_{ij} \in \mathbb{N}^{N \times N}\}$, em que N é o número de nós na rede e x_{ij} é o tráfego, em Gbit/s, do nó i para o nó j .

A RNN usa T passos de tempo para prever a saída Y , onde Y será o passo de tempo $T + 1$. Em outras palavras, T matrizes de tráfego $X^1, X^2, \dots, X^t, \dots, X^T$ são usadas para prever a matriz de tráfego seguinte: $X^{T+1} = Y$.

A RNN com LSTM foi implementada usando a *Application Programming Interface* API do Keras (CHOLLET et al., 2015). Cada matriz de tráfego foi convertida em um vetor unidimensional x_k de largura N^2 , em que $k = i \times N + j$. Para que a entrada se torne T vetores de largura N^2 . Apesar das matrizes terem sido convertidas em vetores, é mais útil pensar na entrada da RNN como uma matriz tridimensional de tamanho $N \times N \times T$.

O modelo de RNN foi designado com 3 camadas recorrentes de T passos de tempo. O T -ésimo passo de tempo da camada 3 gera como saída $N \times N$ características que servem de entrada para uma FFNN completamente conectada que, por sua vez, fornece uma matriz de tráfego prevista X^{T+1} , como saída. A Figura 4.5 ilustra a arquitetura da RNN. Os parâmetros da RNN como w_{aa} e a_0 são compartilhados pelas 3 camadas. Os parâmetros w_{aa} e $a_1, a_2, a_3 \dots a_t$ são passados horizontalmente, para cada camada recorrente. A saída de uma camada recorrente é passada verticalmente ao longo da hierarquia.

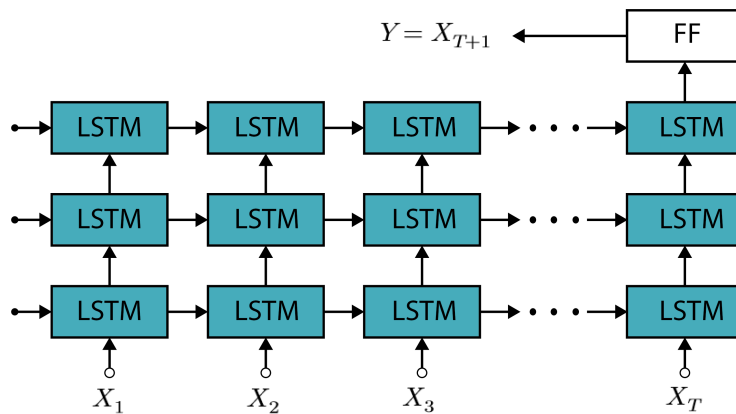


Figura 4.5: Arquitetura da RNN proposta

Para testar a arquitetura proposta, a RNN foi treinada com dados reais da rede Abilene (BALON et al., 2007).

4.3 TREINAMENTO

Os modelos de RNN foram treinados com 8128 exemplos de treinamento, usando o erro médio quadrático ou *Mean Square Error* (MSE) como função de custo. isto é, os valores previstos y^i são comparados com os valores reais $\hat{y}^i$. Por intermédio da Equação (4.11), o erro é elevado ao quadrado e, então a média é feita com todos as amostras usadas como exemplo de treinamento. a RNN é treinada ajustando seus pesos de forma a minimizar a função custo, usando um algoritmo iterativo conhecido como método do gradiente estocástico.

$$\text{MSE} = \frac{1}{2n} \sum_{i=1}^n (y^i - \hat{y}^i)^2. \quad (4.11)$$

Para o método do gradiente estocástico, foram usados lotes de 32 amostras e o treinamento foi executado até a atingir a convergência do MSE, que foi em torno de 0.0019 para a rede Abilene e 0.00048 para a rede GEANT.

Depois do treinamento, os modelos de RNN receberam as duas primeiras amostras de suas respectivas matrizes de séries temporais e preveram a terceira. Então, a RNN usou a segunda e a terceira matriz para prever a quarta e assim sucessivamente, até que a toda a série temporal fosse prevista.

A Figura 4.6 mostra os resultados para a rede GEANT, e a Figura 4.7 mostra os resultados para a rede Abilene. Para ambos os casos, a série temporal colorida de verde é a série original, extraída do banco de dados anonimizado. A curva lilás pontilhada é a série temporal prevista usando o respectivo modelo de RNN.

Para a série temporal prevista, as amostras de 8129 a 10160 não estavam presentes nos dados de treinamento, dessa forma, estas previsões foram baseadas completamente nos padrões aprendidos pela RNN a partir dos dados inseridos previamente na rede durante a fase de treinamento.

Para a rede GEANT, o modelo de RNN fez previsões precisas e foi até capaz de prever oscilações abruptas na demanda de tráfego. Para a rede Abilene, na Figura 4.7, a linha base da previsão é mantida próxima à da série original, mas o modelo tem dificuldades de prever oscilações de tráfego mais drásticas.

De forma geral, é possível ver, por estes segmentos da série temporal que os modelos de RNN são capazes de internalizar padrões percebidos nos dados de treinamento e treinar um modelo generalizado para fazer previsão de tráfego através de regressão.

Os erros de previsão cometidos pelas RNNs podem ser significativos para alocação dinâmica em tempo real, mas podem ser desprezados para aplicações de tráfego incremental, onde apenas a tendência geral do tráfego é importante.

Analiticamente, a performance da RNN pode ser determinada considerando o MSE entre a saída da rede e o resultado esperado (alvo). A Tabela 4.1 mostra o MSE para os dados de treinamento e para os dados previstos (dados de teste).

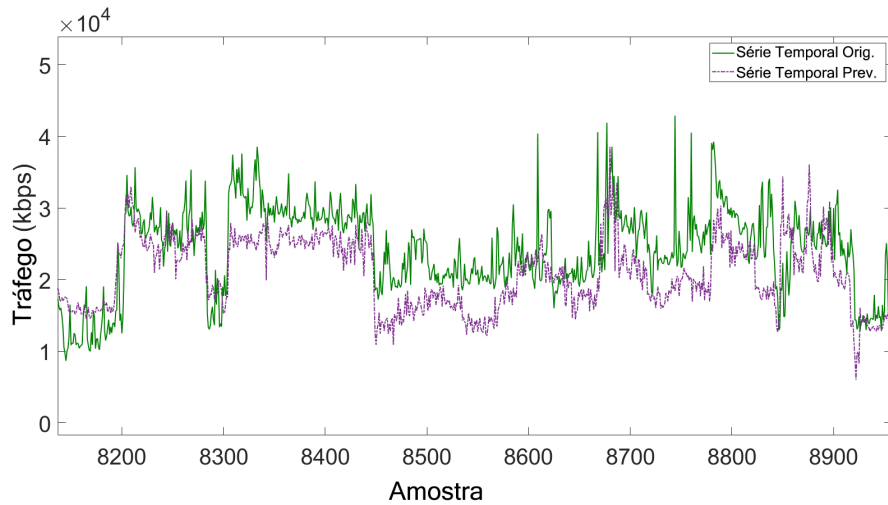


Figura 4.6: Série temporal original e prevista para a demanda de tráfego do um par de nós anônimos origem-destino da rede GEANT.

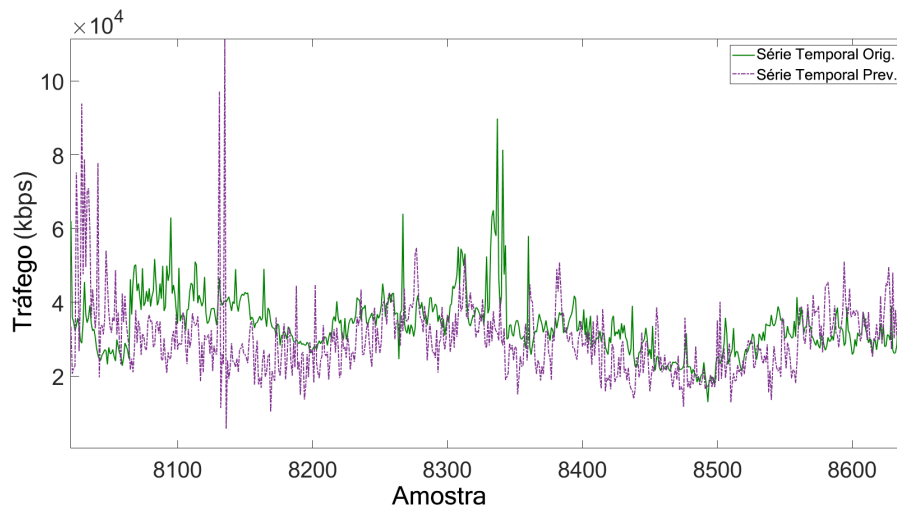


Figura 4.7: Série temporal original e prevista para a demanda de tráfego do um par de nós anônimos origem-destino da rede Abilene.

4.4 RESULTADOS DO RMSA INCREMENTAL

Depois de confirmar a eficiência dos modelos de RNN para previsão de tráfego, um conjunto de ferramentas para previsão instantânea de tráfego foram definidas. O próximo passo, então, é usar as previsões feitas na Seção 4.3 para a estimar a probabilidade bloqueio para a rede Abilene e GEANT, fazendo simulações com valores realistas de tráfego. No entanto, os valores de tráfego contidos no banco de dados são do tráfego instantâneo e não podem ser usados para o caso incremental do RMSA.

Regressão linear foi usada para encontrar a tendência das séries temporais da rede Abilene e GEANT. Como a tendência é crescente, ela pode ser usada para simular tráfego

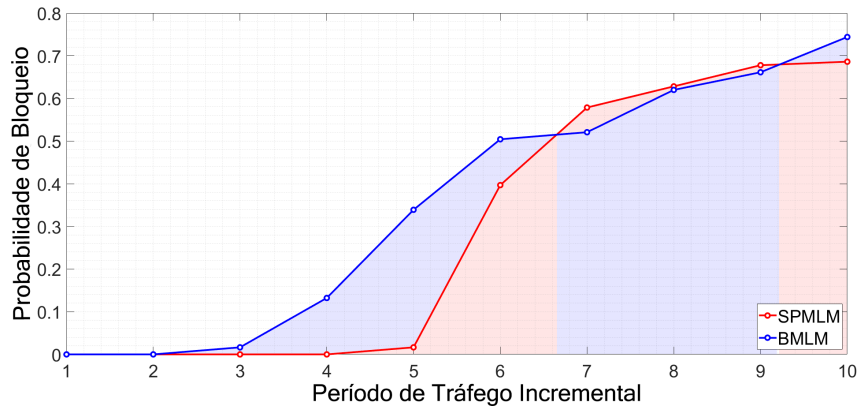
Tabela 4.1: Precisão das RNNs para as previsões de tráfego das redes Abilene e GEANT

Rede	Exemplos de Treinamento	MSE do Treino	MSE do Teste
Abilene	0-8128	0.00193	0.0026
GEANT	0-8128	0.00048	0.00058

incremental, se feita amostragem das séries temporais em intervalos regulares para obter matrizes de tráfego relativas ao período de amostragem. Como a tendência encontrada foi crescentes, a matriz de tráfego amostrada no período t terá maior demanda que a matriz de tráfego amostrada no período $t - 1$, garantindo o incremento de tráfego a cada período. Este processo foi usado para extrair as matrizes de tráfego para ambas as redes e simular o comportamento da probabilidade de bloqueio, assim como na secção anterior.

Repetindo os parâmetros de simulação da Secção 3.3, usando as matrizes de tráfego previstas como para alimentar o Algoritmo 1, e usando os quatro formatos de modulação da Tabela 2.1, foram encontrados padrões similares aos obtidos na Secção 3.3, confirmando os resultados.

A curva de probabilidade de bloqueio, para o SPMLM e o BMLM pode ser vista na Figura 4.4 para a rede Abilene e, para a rede GEANT na Figura 4.4. Para 10 períodos de tráfego incremental. Analisando as curvas, é possível concluir que, em geral, o SPMLM teve melhor performance para todos os períodos de tráfego incremental, exceto durante os períodos de tempo 7, 8 e 9, na rede Abilene.

**Figura 4.8:** Probabilidade de bloqueio para o período de tráfego incremental, para ambas as heurísticas de RMSA, com dados gerados pela RNN para a rede Abilene.

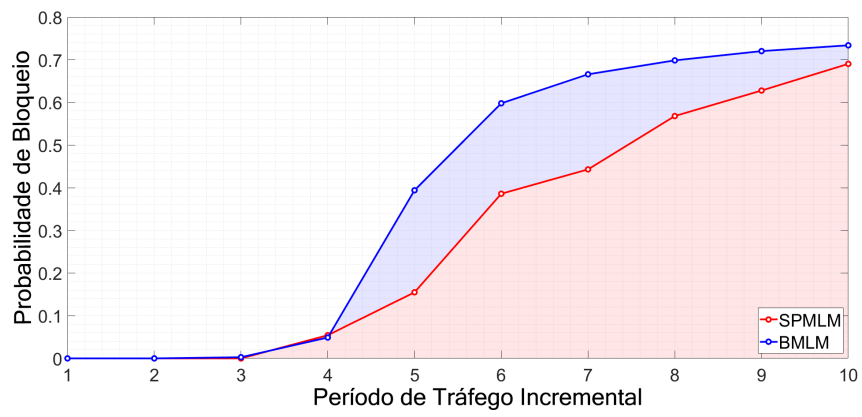


Figura 4.9: Probabilidade de bloqueio para o período de tráfego incremental, para ambas as heurísticas de RMSA, com dados gerados pela RNN para a rede GEANT

CONCLUSÕES

5.1 DISCUSSÃO E CONCLUSÕES FINAIS

Esta dissertação fornece novos métodos multiperíodo para o planejamento de redes ópticas elásticas pela análise da capacidade consumida e da probabilidade de bloqueio ao longo do tempo, em uma abordagem incremental. Uma formulação de MILP foi proposta para resolver o RMSA incremental de forma otimizada. Esta formulação foi testada em relação às duas heurísticas de RMSA incremental, sendo elas uma heurística de roteamento por caminho mais curto (SPMLM) e uma heurística de roteamento por balanceamento de carga (BMLM).

A formulação obteve melhores resultados para a rede pequena em que foi testada, mas a formulação de MILP não é computacionalmente viável, por isso, para redes grandes, apenas as heurísticas foram analisadas, pois podem resolver o problema em um espaço muito mais curto de tempo. A diferença de tempo de processamento é da ordem de dias e não cabe fazer comparações precisas.

Para as simulações com tráfego incremental, primeiramente, foram usadas matrizes com uma demanda de tráfego gerada aleatoriamente e alimentadas às heurísticas, para que a alocação fosse feita, para um total de vinte e três topologias físicas de *backbone* e as condições das redes fossem estudadas na forma da probabilidade de bloqueio.

Quando apenas um formato de modulação podia ser usado, o BMLM mostrou melhor probabilidade de bloqueio média para os primeiros períodos simulados, mas demonstrou pior performance a longo prazo, à medida que a quantidade de conexões aumentava. Este fenômeno é justificado pela exaustão intensiva de recursos e maior fragmentação de espectro causada por este algoritmo. É importante mencionar que uma menor probabilidade de bloqueio inicial por ser útil para provedores de internet, buscando por formas de atrasar a instalação de novos recursos na rede. Nesse contexto, é recomendado o uso do BMLM nos estágios iniciais, e então migrar para o SPMLM, quando se tornar vantajoso.

Para as simulações em que quatro formatos de modulação foram usados, não houve vantagem alguma em usar o BMLM ao invés do SPMLM, exceto para duas redes. A probabilidade de bloqueio se manteve alta durante todos os períodos incrementais, para

o BMLM. A habilidade dos caminhos mais curtos de usar formatos de modulação com maior eficiência espectral fez do SPMLM, uma opção mais vantajosa. Até mesmo nos períodos iniciais, o BMLM não foi capaz de entregar uma performance comparável. No entanto, para a rede ACONET e a rede PIONIER, o BMLM teve menor probabilidade de bloqueio nos dois primeiros períodos de tráfego incremental, mostrando a importância das simulações específicas para a topologia.

Nesta dissertação, também foi proposto um modelo de RNN com LSTM, para fazer previsão de tráfego, baseado em dados reais de fluxo de tráfego para as topologias das redes Abilene e GEANT. Cada modelo foi treinado com 8128 amostras, provindos de um banco de dados anonimizado, até que a rede atingisse a convergência da função de custo, com um MSE de teste de 0,00193 para a rede Abilene e 0,00048 para a GEANT. Estes modelos, então foram usados para prever as matrizes de tráfego, a partir das duas primeiras amostras, e gerar as 10160 amostras seguintes. O MSE da previsão foi de 0,0026 para a rede Abilene e 0,00058 para a rede GEANT, confirmando a viabilidade do uso de RNNs para a previsão de tráfego, já que confirmam os valores publicados em modelos parecidos de RNN para previsão de tráfego.

As simulações propostas para planejamento de redes também foram feitas com os dados gerados pelos modelos de RNN. Para isso, as matrizes de tráfego instantâneo foram convertidas em matrizes de tráfego agregado fazendo a regressão linear das séries temporais relativas a cada par origem-destino. Feito isso, as séries temporais foram amostradas, fornecendo matrizes de tráfego incremental que foram postas à simulação.

Para as simulações com as matrizes previstas, usou-se apenas quatro formatos de modulação disponíveis e os resultados foram análogos aos obtidos com as matrizes geradas aleatoriamente. O SPMLM obteve menor probabilidade de bloqueio para ambas as redes.

De forma geral, o SPMLM parece ser uma melhor opção de heurística para tráfego incremental, caso o uso de formulações de otimização, como o proposto neste trabalho, estejam indisponíveis. A heurística de balanceamento de carga é eficiente em balancear a distribuição de SFs na rede, mas suas desvantagens sobrepõem suas vantagens e o SPMLM, mesmo sendo a heurística mais simples, é melhor para o RMSA incremental, de forma a minimizar a probabilidade de bloqueio para o tráfego incremental.

Os resultados mostraram um padrão consistentes tanto com simulações de tráfego aleatório quanto com simulações com tráfego previsto por RNNs, com dados reais. O trabalho provê informações sobre o comportamento incremental, ao longo do tempo, da probabilidade de bloqueio que pode ajudar administradores de redes a fazer o planejamento multiperíodo de suas redes a partir da probabilidade de bloqueio e escolher o melhor método de RMSA, de forma a consumir menos recursos da rede, o que pode reduzir custos operacionais e aumentar a longevidade da rede projetada.

5.2 TRABALHOS FUTUROS

- A formulação de MILP apresentada nesta dissertação, apesar de útil, pode se tornar inviável de ser usada a depender da rede e do *hardware* disponível. Para resolver este problema, é possível fazer formulações que já comecem com os dados de rotas predeterminadas para as demandas, de forma a considerar menos variáveis no pro-

blema e chegar a um resultado de forma mais rápida. Essa formulação pode ser construída e usada como alternativa para a formulação fornecida aqui.

- As RNNs propostas para fazer a previsão de matrizes de tráfego, são arquiteturas cujo tamanho depende da rede a ser considerada pois o número de entradas da rede e o número de saídas é uma função do número de nós da rede, fazendo com que uma arquitetura tenha que ser projetada para cada rede cuja previsão de tráfego seja desejada (assumindo a disponibilidade de dados). É bastante provável que este fato não possa ser remediado, pois o tráfego de redes maiores têm maior quantidade de correlações entre seus nós, fazendo-se a necessidade de que a RNN, proposta a aprender os padrões do fluxo de tráfego na rede, seja mais complexa, de forma a ter variáveis o suficiente para modelar os padrões e relações do tráfego de forma correta. No entanto, para propostas futuras, seria válido investigar a possibilidade de se fazer uma ANN que seja independente da rede óptica.
- Apesar de fornecer as estratégias para o planejamento multiperíodo de redes ópticas através de métodos matemáticos e heurísticas de RMSA incremental, este trabalho não provê dados dos custos com equipamento e energia necessários para manter a probabilidade de bloqueio dentro de algum padrão de QoS. Essa análise pode ser feita e estimativa de economia financeira pode ser extraída posteriormente.

REFERÊNCIAS BIBLIOGRÁFICAS

- ADAMS, M. *ROADM and Wavelength Selective Switches*. [S.l.]: JDSU, 2008.
- AGRAWAL, G. P. *Fiber-optic communication systems*. [S.l.]: John Wiley & Sons, 2013.
- AIBIN, M. Traffic prediction based on machine learning for elastic optical networks. *Optical Switching and Networking*, Elsevier, v. 30, p. 33–39, 2018.
- AMAR D., R. E. L. B. N. A. J. L. C. P. N. Spectrum fragmentation issue in flexible optical networks: Analysis and good practices. *Photonic Network Communications*, v. 29, n. 3, p. 230–243, 2015.
- AZZOUNI, A.; PUJOLLE, G. Neutm: A neural network-based framework for traffic matrix prediction in sdn. *arXiv preprint arXiv:1710.06799*, 2017.
- AZZOUNI, A.; PUJOLLE, G. Neutm: A neural network-based framework for traffic matrix prediction in sdn. *arXiv preprint arXiv:1710.06799*, 2017.
- BALON, S. et al. Traffic engineering an operational network with the totem toolbox. *IEEE Transactions on Network and Service Management*, IEEE, v. 4, n. 1, p. 51–61, 2007.
- BARRY, R. A.; HUMBLET, P. A. Models of blocking probability in all-optical networks with and without wavelength changers. In: IEEE. *Proceedings of INFOCOM'95*. [S.l.], 1995. v. 2, p. 402–412.
- BENGIO, Y.; BOULANGER-LEWANDOWSKI, N.; PASCANU, R. Advances in optimizing recurrent networks. In: IEEE. *2013 IEEE International Conference on Acoustics, Speech and Signal Processing*. [S.l.], 2013. p. 8624–8628.
- BEPPU, S. et al. 2048 qam (66 gbit/s) single-carrier coherent optical transmission over 150 km with a potential se of 15.3 bit/s/hz. *Optics express*, Optical Society of America, v. 23, n. 4, p. 4960–4969, 2015.
- BISHOP, C. M. *Pattern recognition and machine learning*. [S.l.]: springer, 2006.
- CHANG, R. W. Synthesis of band-limited orthogonal signals for multichannel data transmission. *Bell System Technical Journal*, Wiley Online Library, v. 45, n. 10, p. 1775–1796, 1966.
- CHOLLET, F. et al. Keras: Deep learning library for theano and tensorflow. *URL: https://keras.io/k*, v. 7, n. 8, p. T1, 2015.

CHRISTODOULOPOULOS, K.; TOMKOS, I.; VARVARIGOS, E. A. Elastic bandwidth allocation in flexible ofdm-based optical networks. *Journal of Lightwave Technology*, IEEE, v. 29, n. 9, p. 1354–1366, 2011.

CISCO, V. *The zettabyte era: trends and analysis. Updated (07/06/2017)*. [S.l.]: ed, 2017.

DJORDJEVIC, I. B.; VASIC, B. Orthogonal frequency division multiplexing for high-speed optical transmission. *Optics Express*, Optical Society of America, v. 14, n. 9, p. 3767–3775, 2006.

Eira, A.; Pedro, J.; Pires, J. Optimal multi-period provisioning of fixed and flex-rate modular line interfaces in dwdm networks. *IEEE/OSA Journal of Optical Communications and Networking*, v. 7, n. 4, p. 223–234, April 2015. ISSN 1943-0620.

ELDADA, L. Advances in roadm technologies and subsystems. In: INTERNATIONAL SOCIETY FOR OPTICS AND PHOTONICS. *Photonic Applications in Devices and Communication Systems*. [S.l.], 2005. v. 5970, p. 597022.

FELDMAN, J. M. Beyond attribution theory: Cognitive processes in performance appraisal. *Journal of Applied psychology*, American Psychological Association, v. 66, n. 2, p. 127, 1981.

FENG, H.; SHU, Y. Study on network traffic prediction techniques. In: IEEE. *Proceedings. 2005 International Conference on Wireless Communications, Networking and Mobile Computing, 2005*. [S.l.], 2005. v. 2, p. 1041–1044.

GKAMAS, V.; CHRISTODOULOPOULOS, K.; VARVARIGOS, E. A joint multi-layer planning algorithm for ip over flexible optical networks. *Journal of Lightwave Technology*, IEEE, v. 33, n. 14, p. 2965–2977, 2015.

GONG, L. et al. A two-population based evolutionary approach for optimizing routing, modulation and spectrum assignments (rmsa) in o-ofdm networks. *IEEE Communications letters*, IEEE, v. 16, n. 9, p. 1520–1523, 2012.

HOCHREITER, S. The vanishing gradient problem during learning recurrent neural nets and problem solutions. *International Journal of Uncertainty, Fuzziness and Knowledge-Based Systems*, World Scientific, v. 6, n. 02, p. 107–116, 1998.

ILOG, Inc. *ILOG CPLEX: High-performance software for mathematical programming and optimization*. 2019. See <https://www.ibm.com/analytics/cplex-optimizer>.

IYER, S.; SINGH, S. P. Multiple-period planning of internet protocol-over-elastic optical networks. *Journal of Information and Telecommunication*, Taylor & Francis, v. 3, n. 1, p. 39–56, 2018.

JAYA, T.; GOPINATHAN, E.; RAJENDRAN, V. Comparison of ber performance of various adaptive modulation schemes in ofdm systems. *Indian Journal of Science and Technology*, v. 9, n. 40, p. 1–7, 2016.

- JINNO, M. et al. Distance-adaptive spectrum resource allocation in spectrum-sliced elastic optical path network [topics in optical communications]. *IEEE Communications Magazine*, IEEE, v. 48, n. 8, p. 138–145, 2010.
- JONCIC, M.; HAUPT, M.; FISCHER, U. Standardization proposal for spectral grid for vis wdm applications over si-pof. In: *Proceedings of POF congress*. [S.l.: s.n.], 2012. p. 351–355.
- JONCIC, M.; HAUPT, M.; FISCHER, U. Standardization proposal for spectral grid for vis wdm applications over si-pof. In: *Proceedings of POF congress*. [S.l.: s.n.], 2012. p. 351–355.
- Lardeux, B.; Nace, D. Multiperiod network design with incremental routing. In: . [S.l.: s.n.], 2007. v. 50, n. 1, p. 109–117.
- LELAND, W. E. et al. On the self-similar nature of ethernet traffic. In: ACM. *ACM SIGCOMM computer communication review*. [S.l.], 1993. v. 23, n. 4, p. 183–193.
- LIU, W. et al. Prediction and correction of traffic matrix in an ip backbone network. In: IEEE. *2014 IEEE 33rd International Performance Computing and Communications Conference (IPCCC)*. [S.l.], 2014. p. 1–9.
- LOWERY, A. J.; DU, L. B.; ARMSTRONG, J. Performance of optical ofdm in ultralong-haul wdm lightwave systems. *Journal of Lightwave Technology*, IEEE, v. 25, n. 1, p. 131–138, 2007.
- MINSKY, M.; PAPERT, S. A. *Perceptrons: An introduction to computational geometry*. [S.l.]: MIT press, 2017.
- MOAYEDI, H. Z.; MASNADI-SHIRAZI, M. Arima model for network traffic prediction and anomaly detection. In: IEEE. *2008 International Symposium on Information Technology*. [S.l.], 2008. v. 4, p. 1–6.
- MORALES, F. et al. Virtual network topology adaptability based on data analytics for traffic prediction. *IEEE/OSA Journal of Optical Communications and Networking*, IEEE, v. 9, n. 1, p. A35–A45, 2017.
- MOSIER, R.; CLABAUGH, R. Kineplex, a bandwidth-efficient binary transmission system. *Transactions of the American Institute of Electrical Engineers, Part I: Communication and Electronics*, IEEE, v. 76, n. 6, p. 723–728, 1958.
- OLIVEIRA, H. M.; FONSECA, N. L. da. Routing, spectrum, core and modulation level assignment algorithm for protected sdm optical networks. In: IEEE. *GLOBECOM 2017-2017 IEEE Global Communications Conference*. [S.l.], 2017. p. 1–6.
- PIÓRO, M.; MEDHI, D. *Routing, flow, and capacity design in communication and computer networks*. [S.l.]: Elsevier, 2004.

RAYMOND, E. S. *The Cathedral and the Bazaar*. 1st. ed. Sebastopol, CA, USA: O'Reilly & Associates, Inc., 1999. ISBN 1565927249.

REIXATS, L. N. Design and implementation of low complexity adaptive optical ofdm systems for software-defined transmission in elastic optical networks. Universitat Politècnica de Catalunya, 2014.

ROORDA, P.; COLLINGS, B. Evolution to colorless and directionless roadm architectures. In: OPTICAL SOCIETY OF AMERICA. *National Fiber Optic Engineers Conference*. [S.l.], 2008. p. NWE2.

ROSÁRIO, J. P. G. F. Mb-ofdm metropolitan networks with concatenation of optical add-drop multiplexers. 2014.

ROUTRAY, S. K. et al. Statistical model for link lengths in optical transport networks. *Journal of Optical Communications and Networking*, Optical Society of America, v. 5, n. 7, p. 762–773, 2013.

SAK, H.; SENIOR, A.; BEAUFAYS, F. Long short-term memory recurrent neural network architectures for large scale acoustic modeling. In: *Fifteenth annual conference of the international speech communication association*. [S.l.: s.n.], 2014.

SALTZBERG, B. Performance of an efficient parallel data transmission system. *IEEE Transactions on Communication Technology*, IEEE, v. 15, n. 6, p. 805–811, 1967.

SHADMAND, A.; SHIKH-BAHA EI, M. Multi-user time-frequency downlink scheduling and resource allocation for lte cellular systems. In: IEEE. *2010 IEEE Wireless Communication and Networking Conference*. [S.l.], 2010. p. 1–6.

SOBH, T. et al. *Novel algorithms and techniques in telecommunications, automation and industrial electronics*. [S.l.]: Springer Science & Business Media, 2008.

SOUMPLIS, P. et al. Multi-period planning with actual physical and traffic conditions. *IEEE/OSA Journal of Optical Communications and Networking*, IEEE, v. 10, n. 1, p. A144–A153, 2018.

Straub, S.; Kirstadter, A.; Schupke, D. A. Multi-period planning of wdm-networks: Comparison of incremental and eol approaches. In: *2006 2nd IEEE/IFIP International Conference in Central Asia on Internet*. [S.l.: s.n.], 2006. p. 1–7.

TAKAGI, T. et al. Dynamic routing and frequency slot assignment for elastic optical path networks that adopt distance adaptive modulation. In: IEEE. *2011 Optical Fiber Communication Conference and Exposition and the National Fiber Optic Engineers Conference*. [S.l.], 2011. p. 1–3.

UHLIG, S. et al. Providing public intradomain traffic matrices to the research community. *ACM SIGCOMM Computer Communication Review*, ACM, v. 36, n. 1, p. 83–86, 2006.

- VELASCO, L. et al. Modeling the routing and spectrum allocation problem for flexgrid optical networks. *Photonic Network Communications*, Springer, v. 24, n. 3, p. 177–186, 2012.
- VELASCO, L. et al. On-demand incremental capacity planning in optical transport networks. *IEEE/OSA Journal of Optical Communications and Networking*, IEEE, v. 8, n. 1, p. 11–22, 2016.
- WANG, Y.; CAO, X.; PAN, Y. A study of the routing and spectrum allocation in spectrum-sliced elastic optical path networks. In: IEEE. *2011 Proceedings Ieee Infocom*. [S.l.], 2011. p. 1503–1511.
- WEN, Y.; ZHU, G. Prediction for non-gaussian self-similar traffic with neural network. In: IEEE. *2006 6th World Congress on Intelligent Control and Automation*. [S.l.], 2006. v. 1, p. 4224–4228.
- WINZER, P. J. High-spectral-efficiency optical modulation formats. *Journal of Lightwave Technology*, IEEE, v. 30, n. 24, p. 3824–3835, 2012.
- YEN, J. Finding the k shortest loopless paths in a network. *Management Science*, v. 17, n. 11, p. 712–16, 1971.
- YIN, Y. et al. Fragmentation-aware routing, modulation and spectrum assignment algorithms in elastic optical networks. In: OPTICAL SOCIETY OF AMERICA. *Optical Fiber Communication Conference*. [S.l.], 2013. p. OW3A–5.
- YOO, M.; QIAO, C.; DIXIT, S. Qos performance of optical burst switching in ip-over-wdm networks. *IEEE Journal on selected areas in communications*, IEEE, v. 18, n. 10, p. 2062–2071, 2000.
- ZHAO, J. et al. Towards traffic matrix prediction with lstm recurrent neural networks. *Electronics Letters*, IET, v. 54, n. 9, p. 566–568, 2018.
- ZHAO, J. et al. Towards traffic matrix prediction with lstm recurrent neural networks. *Electronics Letters*, IET, v. 54, n. 9, p. 566–568, 2018.
- ZHAO, J. et al. Distance-adaptive routing and spectrum assignment in ofdm-based flexible transparent optical networks. *Photonic Network Communications*, Springer, v. 27, n. 3, p. 119–127, 2014.
- ZOU, W. Y.; WU, Y. Cofdm: An overview. *IEEE transactions on broadcasting*, IEEE, v. 41, n. 1, p. 1–8, 1995.