



Projeto Brasil 6G Fase III

Atividade 2.2 - Concepção da Rede de Acesso para Rede 6G - Parte 2



Histórico de Atualizações:

Versão	Data	Autor(es)	Notas
1	10/07/2026	Adão Boava Albuquerque Parente Bartolomeu F. Uchôa-Filho Bianca S. de C. da Silva Bruna do Nascimento Benevides Christian Mailer Edgar Eduardo Benitez Olivo Eduardo N. Velloso Felipe Augusto Tondo Francisco Rafael Marques Lima Francisco R. Albuquerque Parente Hirley Alves Igor M. Guerreiro Isaac Andrés Balbuca Guachizaca Joanna Manjarres Jonathan N. Gois José Cândido S. Santos Filho José Ferreira de Rezende Luciano L. Mendes Luiz A. M. Pereira Mariana B. de Mello Matheus D. Carneiro Maykon Renan Pereira da Silva Michel Bernardo de Paiva Miguel Gutiérrez Gaitán Oana Iova Onel L. A. López Osmel M. Rosabal Paulo Cardieri Plínio Santini Dester Raphael Parreira de Souza Richard Demo Souza Samuel Montejo Sánchez Samuel S. Silva Stevan Grubisic Victor F. Monteiro Victoria D. P. Souto Yuri C. B. Silva	Elaboração de conteúdo
2	31/07/2026	Daniely Gomes Silva Filipe Henrique Cardoso José Cândido Silveira Santos Filho Luciano Leonel Mendes Paulo Cardieri Richard Demo Souza	Revisão de texto

Lista de Figuras

1	Modelo do sistema: um <i>Power Beacon</i> (PB) com N antenas, auxiliado por uma <i>Beyond-Diagonal RIS</i> (BD-RIS) com M elementos, alimenta um dispositivo de antena única.	7
2	Energia coletada em função do número de elementos M	9
3	Energia coletada em função do número de antenas N	9
4	Energia coletada vs. fator de Rician κ_D	9
5	Arquitetura da rede neural artificial do tipo multi-layer perceptron utilizada como bloco de <i>Digital Pre-Distortion</i> (DPD).	12
6	Diagrama de blocos do transceptor downlink <i>Faster-than-Nyquist</i> (FTN)- <i>Generalized Frequency Division Multiple Access</i> (GFDMA) integrado ao enlace <i>Analog Radio over Fiber</i> (A-RoF), incluindo o bloco de DPD baseado em ANN na central office/transmissor.	12
7	Desempenho de <i>Bit Error Rate</i> (BER) do sistema FTN- <i>Generalized Frequency Division Multiplexing</i> (GFDMA) transmitido sobre um enlace A-RoF, considerando cenários com e sem DPD.	14
8	Transmissor OFDM com redução de PAPR baseada em RL [1].	18
9	Curvas CCDF do PAPR para Q-Learning, rede neural, referência e PTS [1].	19
10	Desempenho de BER sob canal AWGN para os métodos analisados [1].	20
11	Compressão orientada a tarefas modelada como codificação de fonte indireta com perdas.	21
12	Codificação de fonte distribuída com ruídos aditivos independentes na observação de cada sensor.	22
13	Modelo de sistema para comunicação semântica orientada a tarefas com coleta de energia.	27
14	Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$, no cenário de <i>Evento Raro</i> com $p_s = 0.50$	31
15	Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$, no cenário de <i>Evento Raro</i> com $p_s = 0.90$	32
16	Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$ no cenário de <i>Fonte Rápida</i> com $p_s = 0.50$	33
17	Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$ no cenário de <i>Fonte Rápida</i> com $p_s = 0.90$	34
18	Duração média do transiente E_n em função da probabilidade de acesso p para diversos valores do limiar A e $n = 100$ nós. Os pontos representam os resultados da simulação do modelo original $\{S(t)\}$, e as curvas os resultados analíticos do modelo aproximado $\{K(t)\}$	44
19	AoI médio do 1-persistent <i>Age-Threshold Slotted ALOHA</i> (TSA) para $n = 100$ nós.	44
20	AoI médio em função de ρ com $n = 60$ nós e $p = 2,5/n$	45
21	Perspectiva geral de uma rede de comunicação IoT direta ao satélite.	54

22	Número médio de bytes recebidos com sucesso por volta <i>versus</i> $\mathcal{X}$ para FChirp-FTP, FChirp-CTP, FChirp-ALOHA e FSMA, em comparação com os protocolos originais [2].	58
23	Porcentagem de dispositivos associados a transmissões (a) sem colisões e (b) com colisões não resolvidas <i>versus</i> $\mathcal{X}$ para as estratégias FChirp-FTP, FChirp-ALOHA, FSMA e FTP.	58
24	Arquitetura <i>Cell-Free</i> (CF)- <i>Massive Multiple-Input Multiple-Output</i> (mMIMO) com múltiplas <i>Central Processing Unit</i> (CPU)s em sua versão escalável.	61
25	Exemplo de cenário utilizado no trabalho. As linhas indicam a associação entre cada CPU e os <i>Access Point</i> (AP)s sob sua responsabilidade. O critério de associação é baseado na menor distância considerando <i>wrap-around</i>	62
26	Comparação de SE e JFI para o método de DCC escalável.	69
27	Alocação de potência segundo o método proposto e o método denominado como CPU-Avg.	71
28	Performance comparison of ML models by isolated propagation type using multimodal context (LiDAR + GPS) with Domain Source $\mathcal{D}_S = s008$ and Domain Target $\mathcal{D}_T = s009$	76
29	Comparação de desempenho entre domínios sob o protocolo de troca de domínios <i>Domain Swap</i> , utilizando contexto multimodal (LiDAR + GPS), com domínio de origem $\mathcal{D}_S = s009$ e domínio de destino $\mathcal{D}_T = s008$	77
30	Distribuição de classes dos índices de feixe para o conjunto de dados s008.	78
31	Distribuição de classes dos índices de feixe para os conjuntos de dados s009.	79
32	Diagrama de Venn mostrando a sobreposição das classes de índice de feixe nos conjuntos de dados s008 e s009 sob condição de propagação.	80
33	Distribuição geográfica dos receptores (Rx) nos cenários s008 e s009 sob propagação LoS. O cenário $s009_{LoS}$ (à direita) atua como um "zoom" espacial de $s008_{LoS}$ (à esquerda), resultando em uma grave deficiência de suporte. As caixas tracejadas em vermelho destacam as regiões superior ($y > 579$) e inferior ($y < 510$) não representadas em $s009_{LoS}$, o que explica o colapso catastrófico do suporte para classes fora do vocabulário ($ I_{overlap} = 41$ de um total de 81 classes ativas).	86
34	Mapa de calor da correlação de postos de Kendall (τ) detalhando as relações entre métricas de divergência geométricas/teórico-informacionais e a acurácia empírica fora da distribuição dos modelos avaliados.	87
35	Associação entre APs e <i>User Equipments</i> (UEs) e alocação de <i>Resource Blocks</i> (RBs) ao longo do tempo em uma rede centrada no usuário com <i>Distributed Multi-input Multi-output</i> (D-MIMO) e <i>Dynamic TDD</i> (D-TDD).	98
36	Resultados de desempenho para as arquiteturas avaliadas.	99
37	Número médio de APs conectados a cada UE para valores de δ	107
38	Impacto do limiar do <i>Dynamic Cooperation Clustering</i> (DCC) (δ) nas métricas de desempenho em <i>Uplink</i> (UL).	107
39	Impacto do limiar do DCC (δ) nas métricas de desempenho em <i>Downlink</i> (DL).	108
40	Comparação de esquemas de controle de potência e agrupamento nas métricas de desempenho de UL.	109
41	Modelo do sistema de satélite <i>Low Earth Orbit</i> (LEO).	113

42	Taxa total de dados alcançada pelos usuários em função de D , para avaliar o desempenho de <i>Multiple Access</i> (MA). A simulação foi conduzida com $\Delta D = 10$ km, considerando $\tau = 0$ para <i>Channel State Information at the Transmitter</i> (CSIT) perfeito e $\tau = 0.05$ para CSIT imperfeito.	120
43	Taxa total de dados alcançada pelos usuários em função de τ , para avaliar o desempenho de MA. A simulação foi conduzida com $\Delta D = 10$ km e $D = 100$ km.	121
44	Modelo do sistema.	123
45	Probabilidade de indisponibilidade do pior usuário em função da <i>Signal-to-Noise Ratio</i> (SNR) de transmissão para cenários homogêneos e heterogêneos com $N_t = 4$, $K = 4$ e $Q = 2$, comparando a alocação de potência otimizada via Progressão Geométrica (PG) com uma abordagem arbitrária.	131
46	Probabilidade de indisponibilidade do pior usuário para $N_t = 4$ e $N_t = 8$ antenas, com $K = 4$ e $Q = 2$, evidenciando o ganho de diversidade espacial e o ganho adicional proporcionado pela otimização via PG.	131
47	Probabilidade de indisponibilidade do pior usuário para diferentes configurações de agrupamento: 4 grupos com 2 usuários ($Q = 4$, $K_q = 2$) versus 2 grupos com 4 usuários ($Q = 2$, $K_q = 4$), com $N_t = 8$ e $K = 8$	132
48	Probabilidade de indisponibilidade do pior usuário para diferentes números de usuários $K = 4, 6$ e 8 , com $N_t = 8$ e $Q = 2$, resultando em tamanhos de grupo $K_q = 2, 3$ e 4 , respectivamente.	132
49	Modelo do sistema <i>Non-Terrestrial Networks</i> (NTN) baseado em NOMA.	137
50	a) Probabilidade de Indisponibilidade (outage) <i>versus</i> a SNR para diferentes valores de α . b) Probabilidade de Indisponibilidade (outage) <i>versus</i> a SNR para diferentes valores de M	140
51	a) Probabilidade de Indisponibilidade (outage) <i>versus</i> a SNR para diferentes valores de correlação. b) Probabilidade de Indisponibilidade (outage) <i>versus</i> a SNR para diferentes valores de σ	141
52	Probabilidade de Indisponibilidade (outage) <i>versus</i> a SNR para diferentes valores de k	141
53	Modelo do sistema NTN.	144
54	(a) Curva de convergência por iteração. (b) Comparação da taxa de transmissão com alocação de potência fixa e PSO. (c) Região de taxas considerando os limites de NOMA e OMA.	150
55	Probabilidade de outage em função da SNR média para diferentes configurações dos parâmetros ρ , M , σ , k e α , incluindo a curva correspondente ao método de otimização proposto.	152

Lista de Tabelas

1	Parâmetros do sistema FTN-GFDM utilizados na avaliação do DPD.	13
2	Modelo de canal utilizado na avaliação do DPD.	13
3	Parâmetros de Simulação do Cenário DtS-IoT com FreeChirp.	57
4	Parâmetros de simulação.	68
5	Comparação de desempenho usando precodificador <i>Maximum Ratio</i> (MR) e $\tau_p = 5$ (com contaminação de pilotos). Parâmetro DCC escalável $\alpha = 0,9$. Parâmetros do DCC CPU- <i>aware</i> : $\gamma_{\text{cpu}} = 0,1$, $\tau_{\text{top}} = 0,3$	70
6	Estatísticas dos <i>dataset</i> s008 e s009.	74
7	Parâmetros usados pelas referencias B. Salehi [3], J. Ruseckas [4] and J. Manjarres[5] nos modelos multimodais	75
8	Comparação entre as técnicas do estado da arte para métricas-chave de entrada multimodal	75
9	Métricas geométricas e de transferibilidade entre os domínios s008 e s009.	84
10	Parâmetros de simulação do sistema com D-MIMO e D-TDD.	97
11	Parâmetros de simulação da rede <i>cell-free</i> com D-TDD	106
12	Parâmetros de simulação da rede LEO-terrestre.	118
13	Parâmetros de simulação do sistema NTN baseado em NOMA.	139

Acrônimos

2G	Segunda Geração de Rede Móvel Celular
3GPP	<i>3rd Generation Partnership Project</i>
4G	Quarta Geração de Rede Móvel Celular
5G	Quinta Geração de Rede Móvel Celular
6G	Sexta Geração de Rede Móvel Celular
A2G	<i>Air-to-Ground</i>
AoI	<i>Age of Information</i>
AP	<i>Access Point</i>
AWGN	<i>Additive White Gaussian Noise</i>
A-RoF	<i>Analog Radio over Fiber</i>
ANN	<i>Artificial Neural Network</i>
BD-RIS	<i>Beyond-Diagonal RIS</i>
BER	<i>Bit Error Rate</i>
BS	<i>Base Station</i>
BPSK	<i>Binary Phase Shift Keying</i>
CAD	<i>Channel Activity Detection</i>
CF	<i>Cell-Free</i>
CLI	<i>Cross Link Interference</i>
CP	<i>Cyclic Prefix</i>
CPU	<i>Central Processing Unit</i>
CRC	<i>Cyclic Redundancy Check</i>
CSI	<i>Channel State Information</i>
CSIT	<i>Channel State Information at the Transmitter</i>
CTP	<i>Controlled Transmit Power</i>
CEO	<i>Chief Executive Officer</i>
CP	<i>Cyclic Prefix</i>
CO	<i>Central Office</i>
CCDF	<i>Complementary Cumulative Distribution Function</i>
CLT	<i>Central Limit Theorem</i>
DCC	<i>Dynamic Cooperation Clustering</i>
DL	<i>Downlink</i>
D-MIMO	<i>Distributed Multi-input Multi-output</i>
DNS	<i>Domain Name System</i>
DoS	<i>Denial of Service</i>

DPD *Digital Pre-Distortion*
D-TDD *Dynamic TDD*
DtS-IoT *Direct-to-Satellite IoT*
DQN *Deep Q-Network*
eCDF *Empirical Cumulative Distribution Function*
EE *Energy Efficiency*
eRAC *Enhanced Remote Areas Communication*
FBS *Fake Base Station*
FFT *Fast Fourier Transform*
FPC *Fractional Power Control*
FSMA *Free Signal Multiple Access*
FTP *Fixed Transmit Power*
FUE *Fake User Equipment*
FTN *Faster-than-Nyquist*
FSPL *Free Space Path Loss*
GFDM *Generalized Frequency Division Multiplexing*
GFDMA *Generalized Frequency Division Multiple Access*
HRSMA *Hierarchical Rate-Splitting Multiple Access*
ICI *Intercarrier Interference*
IFFT *Inverse Fast Fourier Transform*
IMSI *International Mobile Subscriber Identity*
IoT *Internet of Things*
IP *Internet Protocol*
ISI *Intersymbol Interference*
IUI *Inter-User Interference*
ISM *Industrial, Scientific and Medical*
KKT *Karush–Kuhn–Tucker*
LEO *Low Earth Orbit*
LOS *Line of Sight*
LSF *Large Scale Fading*
LTE *Long-Term Evolution*
MA *Multiple Access*
MICP *Mixed-Integer Convex Optimization Problem*
MIMO *Multiple-Input Multiple-Output*
MINLP *Mixed-Integer Nonlinear Problem*

MISO *Multiple-Input Single-Output*
MiTM *Man-in-The-Middle*
MLP *Multi-Layer Perceptron*
mMIMO *Massive Multiple-Input Multiple-Output*
MMSE *Minimum Mean Square Error*
MP *Max Power*
MR *Maximum Ratio*
MRT *Maximal Ratio Transmission*
MSE *Mean Squared Error*
MTC *Machine Type Communications*
MZM *Mach-Zehnder Modulator*
ML *Machine Learning*
NAS *Non-Access Stratum*
NLOS *Non-Line of Sight*
NOMA *Non-Orthogonal Multiple Access*
NTN *Non-Terrestrial Networks*
NR *New Radio*
OFDM *Orthogonal Frequency Division Multiplexing*
OMA *Orthogonal Multiple Access*
OOBE *Out-of-Band Emissions*
OP *Outage Probability*
PAoI *Peak Age of Information*
PAPR *Peak-to-Average Power Ratio*
PB *Power Beacon*
PDF *Probability Density Function*
PSD *Power Spectral Density*
PSO *Particle Swarm Optimization*
PTS *Partial Transmit Sequence*
PWS *Public Warning System*
PG *Progressão Geométrica*
QAM *Quadrature Amplitude Modulation*
QoS *Quality of Service*
RB *Resource Block*
ReLU *Rectified Linear Unit*
RF *Radiofrequência*

RIS *Reconfigurable Intelligent Surfaces*
RRC *Radio Resource Control*
RRM *Radio Resource Management*
RSMA *Rate-Splitting Multiple Access*
RL *Reinforcement Learning*
RZF *Regularized Zero-Forcing*
SAT *Stationary Age-based Thinning*
SC *Superposition Coding*
SCA *Successive Convex Approximation*
SDR *Software Defined Radio*
SDMA *Space Division Multiple Access*
SE *Spectral Efficiency*
SELU *Scaled Exponential Linear Unit*
SIC *Successive Interference Cancellation*
SINR *Signal to Interference plus Noise Ratio*
SISO *Single-Input Single-Output*
SLM *Selected Mapping*
SMS *Short Message Service*
SNR *Signal-to-Noise Ratio*
SUCI *Subscription Concealed Identifier*
TDD *Time Division Duplex*
TDMA *Time Division Multiple Access*
TMSI *Temporary Mobile Subscriber Identity*
TS *Technical Specification*
TSA *Age-Threshold Slotted ALOHA*
TVWS *TV White Space*
U *Unlicensed*
UAV *Unmanned Aerial Vehicle*
UE *User Equipment*
UHF *Ultra High Frequency*
UL *Uplink*
ULA *Uniform Linear Array*
VHF *Very High Frequency*
V2X *Vehicle-to-Everything*
ZF *Zero Forcing*
WET *Wireless Energy Transfer*
WF *Water-Filling*

Índice

1	Introdução Geral	1
2	Projeto Conjunto do <i>Beamforming</i> Analógico e da BD-RIS para Transferência Eficiente de Energia Sem Fio	5
2.1	Introdução	5
2.2	Modelo do Sistema	6
2.2.1	Arquiteturas de BD-RIS	7
2.3	Resultados de Simulação	8
2.4	Conclusões	9
3	Análise de Pré-Distorção Digital em Sistemas FTN-GFDMA Integrados ao A-RoF	10
3.1	Introdução	10
3.2	Digital Pre-Distortion Baseado em Redes Neurais	11
3.3	Resultados e Discussões	13
3.4	Conclusão	15
4	Aprendizado por Reforço para Reduzir PAPR em Sistemas de Comunicação	16
4.1	Introdução	16
4.2	Desenvolvimento	16
4.2.1	Fundamentos do Sistema OFDM	17
4.2.2	Caracterização do PAPR	17
4.2.3	Fundamentos do Q-Learning	17
4.2.4	Aplicação do Q-Learning à Redução de PAPR	18
4.3	Resultados Numéricos	18
4.3.1	Análise de PAPR	18
4.3.2	Análise de BER	18
4.4	Conclusão	19
5	Codificação Multiterminal para Fontes Gaussianas com Ruído Não-Estacionário	21
5.1	Introdução	21
5.2	Formulação do Problema	22
5.3	Resultados	24
5.4	Conclusão	25
6	Rastreamento Remoto de Processos Baseado em Comunicação Orientada a Tarefas e Predição	26
6.1	Introdução	26
6.2	Modelo de Sistema	27
6.3	Políticas de Transmissão	28
6.4	Resultados Numéricos	29
6.4.1	Configuração das Simulações e Cenários	29
6.4.2	Métricas de Desempenho	30
6.4.3	Resultados para o Cenário de Evento Raro	31
6.4.4	Resultados para o Cenário de Fonte Rápida	32
6.4.5	Discussão	33

6.5	Conclusões e Trabalhos Futuros	34
7	Novo Protocolo de Acesso Aleatório baseado em Age of Information	36
7.1	Introdução	36
7.2	Modelo do Sistema	37
7.3	Resultados Limites	38
7.4	Análise do Transiente	42
7.5	Simulações	44
7.6	Conclusão	46
8	Fake Base Station: Ameaças, Detecção e Mitigação.	47
8.1	Introdução	47
8.2	Métodos de Ataque Baseados em FBS	48
8.3	Tipos de Ataque	49
8.4	Detecção e Mitigação	50
8.4.1	Detecção	50
8.4.2	Mitigação	50
8.5	Boas Práticas de Segurança	51
8.6	Correções e Recomendações para o 6G	51
8.7	Conclusão	52
9	Estratégias de Uplink NOMA Baseadas em FreeChirp para Redes IoT de Comunicação Direta ao Satélite	53
9.1	Introdução	53
9.2	Metodologia	54
9.3	Resultados	57
9.4	Conclusão	59
10	Otimização Conjunta Escalável de Recursos para Sistemas Cell-Free Massive MIMO Multi-CPU	60
10.1	Introdução	60
10.2	Modelo do Sistema	62
10.2.1	Topologia da Rede e Arquitetura Multi-CPU	62
10.2.2	Modelo de Canal	63
10.2.3	Dynamic Cooperation Clustering (DCC)	63
10.2.4	Transmissão em <i>downlink</i>	64
10.2.5	Problema de Otimização	65
10.3	Métodos Propostos	66
10.3.1	DCC Escalável	66
10.3.2	DCC Escalável com CPU- <i>aware</i>	67
10.3.3	Alocação de Potência Distribuída	67
10.4	Resultados	68
10.5	Conclusão	69
11	Análise de Generalização e Divergência de Domínio em Seleção de Feixe Baseada em Informações de Contexto	72
11.1	Introdução	72
11.2	Desenvolvimento das Atividades	73

11.2.1	Configuração experimental multimodal	73
11.2.2	Arquiteturas de <i>baseline</i> Multimodais Avaliadas	74
11.2.3	Avaliação Empírica e a Barreira da Generalização entre Domínios	75
11.2.4	Caracterização Estatística dos Datasets	77
11.2.5	Análise de Divergência entre Domínios	80
11.3	Principais Resultados Obtidos	86
11.4	Conclusão e trabalhos futuros	88
12	Gerenciamento de Recursos de Rádio em Redes MIMO Distribuídas Centra-	
	das no Usuário com TDD Dinâmico	90
12.1	Introdução	90
12.2	Modelo do Sistema	90
12.3	Formulação do Problema	92
12.4	Reformulação Proposta	93
12.5	Soluções de Referência	95
12.5.1	Remoção da Operação Centrada no Usuário com D-MIMO	96
12.5.2	Remoção do D-TDD (TDD Estático)	96
12.5.3	Configurações de Referência	96
12.6	Avaliação de Desempenho	97
12.7	Conclusões e Trabalhos Futuros	100
13	Agrupamento Centrado no Usuário e Controle de Potência em Redes Cell-	
	Free com TDD Dinâmico	101
13.1	Introdução	101
13.2	Modelo de Sistema	102
13.2.1	Modelo de Propagação	102
13.2.2	Combinadores	103
13.2.3	Modelo de Sinal	103
13.3	Controle de Potência	104
13.4	Estratégia de Agrupamento em <i>Dynamic TDD</i> (D-TDD)	105
13.5	Resultados de Simulação	105
13.5.1	Análise do Limiar <i>Dynamic Cooperation Clustering</i> (DCC)	106
13.5.2	Análise do Controle de Potência e Agrupamento	109
13.6	Conclusões	110
14	Estratégias de Alocação de Potência para Mensagens Comuns e Privadas em	
	Redes RSMA LEO-Terrestres	111
14.1	Introdução	111
14.2	Modelo do Sistema e formulação do problema	111
14.2.1	Modelo do Sistema	112
14.2.2	Modelo de Canal e Informação de Estado	113
14.3	Modelo de Transmissão com RSMA	113
14.3.1	Formulação do Problema de Otimização	114
14.4	Solução Proposta	115
14.4.1	Análise dos Precodificadores	115
14.4.2	Divisão de Potência	116
14.5	Avaliação de Desempenho	118
14.5.1	Detalhes de Implementação	118

14.5.2 Soluções de Referência	118
14.6 Resultados	119
14.7 Conclusão	121
15 Avaliação de RSMA Hierárquico em Sistemas Downlink com Múltiplas Antenas: Outage e Alocação de Potência	122
15.1 Introdução	122
15.2 Modelo do Sistema	123
15.2.1 Decodificação dos Sinais e Análise da SINR	126
15.3 Análise da Probabilidade de Indisponibilidade	127
15.4 Alocação de Potência Min–Max Baseada em Programação Geométrica	128
15.4.1 Formulação do Problema	128
15.4.2 Reformulação via Programação Geométrica	129
15.5 Resultados Numéricos	129
15.6 Conclusão	132
16 Desempenho de Sistemas NTN Baseados em NOMA sob Desvanecimento Generalizado	134
16.1 Introdução	134
16.2 Modelo de Canal de Desvanecimento	134
16.3 Modelo do Sistema	136
16.3.1 <i>Superposition Coding</i>	136
16.3.2 <i>Successive Interference Cancellation</i>	136
16.3.3 Imperfeições do SIC e degradação de <i>hardware</i>	137
16.3.4 Geometria do modelo e caracterização do canal	137
16.4 Resultados Numéricos e Discussão	138
16.5 Conclusão	141
17 Otimização de Redes 6G NTN-NOMA em Canais de Desvanecimento Generalizado	143
17.1 Introdução	143
17.2 Modelo do Sistema	143
17.2.1 Modelo de Desvanecimento	144
17.2.2 Geometria do UAV e Caracterização do Canal	145
17.3 Transmissão NOMA e Codificação por Superposição	146
17.4 Estratégia de Otimização da Alocação de Potência	147
17.4.1 Formulação do Problema	147
17.4.2 Algoritmo de Otimização por Enxame de Partículas	148
17.5 Resultados e Discussões	149
17.6 Conclusão	153
18 Conclusão Geral	154

1 Introdução Geral

Esse relatório apresenta resultados de estudos exploratórios sobre tecnologias inovadoras e habilitadoras das redes 6G, conduzidos durante o segundo ano da Atividade 2.2, do Projeto Brasil 6G - Fase III. Os estudos concentraram-se em aspectos da rede de acesso, abrangendo as camadas física e de controle de acesso.

O relatório é composto por dezessete capítulos (além dos capítulos Introdução e Conclusões Gerais), cada um dedicado aos temas investigados pelas equipes das instituições participantes da Atividade 2.2: Instituto Nacional de Telecomunicações, Universidade Estadual de Campinas, Universidade Federal de Santa Catarina, Universidade Federal do Rio de Janeiro, Universidade Federal do Ceará e Fundação CPqD. A seguir, apresenta-se uma breve descrição do conteúdo de cada capítulo.

No Capítulo 2, investiga-se *Wireless Energy Transfer* (WET) por Radiofrequência (RF) para redes *Internet of Things* (IoT) de Sexta Geração de Rede Móvel Celular (6G), integrando um *Power Beacon* (PB) a uma *Beyond-Diagonal RIS* (BD-RIS), arquitetura que controla conjuntamente fase e magnitude do sinal e supera as *Reconfigurable Intelligent Surfaces* (RIS) convencionais. Propõe-se uma solução de *beamforming* analógico baseada em *Particle Swarm Optimization* (PSO) para maximizar a energia coletada, considerando arquiteturas de conexão única, totalmente conectada e conectada por grupos. Os resultados mostram ganhos de ao menos 35% (totalmente conectada) e até 30% (conectada por grupos) frente à RIS clássica, evidenciando o potencial da BD-RIS para WET.

O Capítulo 3 aborda a integração da técnica FTN-GFDMA em sistemas A-RoF (Analog Radio over Fiber), visando principalmente a cobertura de áreas remotas (eRAC). Como o grande desafio prático são as distorções não lineares geradas pela cadeia óptica (especialmente pelo modulador Mach-Zehnder) somadas à interferência do próprio FTN, o texto foca em uma solução bem interessante: o uso de Pré-Distorção Digital (DPD) baseada em redes neurais artificiais. O detalhe importante é que esse bloco de DPD fica na Central Office, mantendo a baixa complexidade das unidades remotas. Os resultados mostraram que a rede neural consegue compensar essas distorções de forma muito eficaz, melhorando consideravelmente o BER do sistema.

No Capítulo 4, aplica-se *Reinforcement Learning* (RL) para reduzir a *Peak-to-Average Power Ratio* (PAPR) em sistemas *Orthogonal Frequency Division Multiplexing* (OFDM), característica que exige amplificadores lineares e limita o desempenho, sobretudo em faixas de potência restrita como *TV White Space* (TVWS). Adota-se o algoritmo *Q-Learning*, que aprende, por interação direta com o ambiente e sem necessidade de bases de dados, combinações de subportadoras piloto que minimizam a PAPR. O método é comparado a uma rede neural supervisionada e à técnica *Partial Transmit Sequence* (PTS). Os resultados indicam que o *Q-Learning* reduz a PAPR com baixa complexidade computacional e leve degradação de *Bit Error Rate* (BER).

No Capítulo 5, aborda-se a codificação de fonte multiterminal, o problema do *Chief Executive Officer* (CEO), para fontes Gaussianas com ruído de observação não-estacionário, no contexto da comunicação orientada a tarefas em 6G. Considerando fontes com estrutura Toeplitz/Sobolev e ruídos com variância suavemente variável no tempo, deriva-se uma caracterização de primeira ordem da função taxa-distorção no domínio tempo-frequência. Demonstra-se que o problema vetorial, inerentemente não-comutativo, reduz-se assintoticamente a subproblemas escalares do CEO acoplados por uma restrição global de distorção, admitindo uma solução do tipo *waterfilling* no plano tempo-frequência, conectando teoria de operadores e teoria da informação clássica.

No Capítulo 6, apresenta-se uma estratégia de comunicação orientada a tarefas para o problema de rastreamento remoto de processos sob restrições de energia, com foco na redução do erro de reconstrução e do consumo energético em sistemas IoT. A principal contribuição consiste na proposta de uma política preditiva de transmissão que combina o conhecimento do estado reconstruído no receptor com a predição da evolução futura de uma fonte markoviana, evitando transmissões desnecessárias quando há expectativa de ressincronização natural. Os resultados obtidos demonstram que, em cenários com severas restrições energéticas, a abordagem proposta preserva energia, reduz o erro médio de reconstrução em relação à política semântica convencional e estabelece uma base promissora para futuras políticas ótimas baseadas em processos de decisão de Markov.

No Capítulo 7, é apresentado um novo protocolo de acesso aleatório baseado na métrica *Age of Information* (AoI) para redes 6G voltadas a aplicações de IoT e *Machine Type Communications* (MTC). A proposta, denominada 1-persistent Age-Threshold Slotted ALOHA (1-persistent TSA), estabelece que cada dispositivo transmita com probabilidade unitária quando a idade da informação atinge um limiar predefinido, adotando retransmissões com probabilidade fixa após falhas de transmissão. O capítulo apresenta a análise do regime estacionário e transiente do protocolo, obtém expressões fechadas para throughput, AoI e *Peak Age of Information* (PAoI) e valida os resultados por meio de simulações computacionais.

O Capítulo 8 trata da ameaça persistente de *Fake Base Stations* (FBS) em redes móveis, desde o 2G até propostas para o 6G. Classificam-se os métodos de ataque (*Man-in-The-Middle*, injeção de mensagens, personificação, *access reject*, *handover* malicioso, *International Mobile Subscriber Identity catching*, escuta passiva) e os tipos de dano resultantes (DoS, *downgrade* tecnológico, espionagem, fraude, dreno de energia). Discutem-se estratégias de detecção baseadas em UE, rede, dispositivos estacionários e abordagens híbridas, além de mitigações como roteamento por reputação. O capítulo propõe seis recomendações estruturais para que o 6G não herde essas vulnerabilidades: autenticação de sinalização, pré-autenticação, integridade nativa do plano de usuário e detecção embutida na arquitetura.

No Capítulo 9, apresentam-se novas estratégias de acesso ao meio para redes IoT com comunicação direta ao satélite, integrando o mecanismo de controle FreeChirp a esquemas *Non-Orthogonal Multiple Access* (NOMA) no domínio da potência. São propostas duas abordagens, denominadas FChirp-*Fixed Transmit Power* (FTP) e FChirp-*Controlled Transmit Power* (CTP), que utilizam um sinal de controle transmitido pelo satélite para coordenar o acesso ao canal e reduzir colisões, preservando a simplicidade operacional e dispensando sincronização rigorosa entre os dispositivos. Os resultados, obtidos a partir de simulações com dados orbitais reais do satélite LacunaSat-3, demonstram ganhos expressivos de capacidade (*goodput*) em cenários de alta densidade de dispositivos, evidenciando o potencial das estratégias propostas para ampliar a escalabilidade e a eficiência de redes *Direct-to-Satellite IoT* (DtS-IoT) de próxima geração.

No Capítulo 10, aborda-se o desafio de escalabilidade em sistemas *Cell-Free Massive MIMO* operando com múltiplas *Central Processing Unit* (CPU). Propõe-se um método de *Dynamic Cooperation Clustering* (DCC) escalável baseado em *K-means* sobre características de larga escala (ganho de canal e distância AP-UE), com uma variante CPU-aware que reduz o número de CPUs envolvidas por usuário. Complementarmente, apresenta-se uma técnica de refinamento de potência (*α -refinement*) de baixa dimensão para melhorar a justiça max-min sem exigir cooperação intensiva entre CPUs. Resultados de simulação mostram ganhos de eficiência espectral e de justiça (Índice de Justiça de *Jam*) em relação a abordagens baseline, mesmo sob contaminação de pilotos, com baixo overhead de sinalização.

No Capítulo 11, é apresentada uma investigação sobre a capacidade de generalização de modelos de aprendizado de máquina aplicados à seleção de feixes baseada em informações de contexto para redes 6G. Para isso, é proposta uma metodologia de avaliação entre cenários, considerando diferentes condições de propagação e mudanças de domínio, com o objetivo de analisar como características estatísticas, geométricas e estruturais dos conjuntos de dados afetam o desempenho dos modelos. Também são empregadas métricas para caracterizar as diferenças entre distribuições e avaliar sua relação com a generalização entre cenários. Os resultados obtidos permitem identificar fatores associados à degradação de desempenho sob mudanças de domínio.

No Capítulo 12, o foco é o gerenciamento de recursos em redes MIMO distribuídas (Cell-Free) centradas no usuário e que operam com TDD Dinâmico. A ideia é lidar com o problema do tráfego assimétrico de UL e DL. O capítulo propõe otimizar de forma conjunta a associação entre APs e usuários, o controle de potência e o escalonamento. Para resolver isso computacionalmente, os autores pegam um problema não linear complexo e aplicam técnicas de linearização e agrupamento para transformá-lo em um problema convexo (MICP). No fim, as simulações comprovam que unir a abordagem centrada no usuário com o D-TDD reduz a probabilidade de outage e melhora a taxa total da rede quando comparado ao TDD estático.

No Capítulo 13, são avaliados os efeitos conjuntos de um algoritmo de reagrupamento dinâmico baseado em DCC e do esquema *Fractional Power Control* (FPC) em uma rede MIMO cell-free operando com D-TDD. O estudo considera métricas de atraso de pacotes, eficiência espectral e eficiência energética, analisando o impacto do reagrupamento dinâmico e do controle de potência no desempenho do sistema. Para isso, é desenvolvido um modelo para redes cell-free com D-TDD e são realizadas simulações para comparar as estratégias propostas com abordagens convencionais de agrupamento e controle de potência.

O Capítulo 14 explora a otimização da alocação de potência em redes de satélites LEO utilizando o acesso múltiplo RSMA (*Rate-Splitting Multiple Access*). O texto mostra como dividir as mensagens em partes comuns e privadas ajuda a gerenciar melhor a interferência, especialmente quando a informação de estado do canal no transmissor (CSIT) é imperfeita. Em vez de usar buscas exaustivas que pesam muito computacionalmente, o estudo traz uma estratégia mais simples usando o método Zero Forcing para a parte privada e um critério max-min para a comum. Os resultados confirmam que essa abordagem do RSMA é bem mais robusta e eficiente que o OMA e SDMA convencionais nesses cenários não terrestres.

No Capítulo 15, apresenta-se um arcabouço teórico para análise e otimização do *Hierarchical Rate-Splitting Multiple Access* (HRSMA) em sistemas downlink com múltiplas antenas sob desvanecimento Nakagami- m . O trabalho obtém expressões analíticas exatas e assintóticas para a probabilidade de indisponibilidade (*outage*), propondo ainda uma estratégia de alocação de potência do tipo min-max, baseada em programação geométrica, para minimizar a pior probabilidade de indisponibilidade entre os usuários. Os resultados analíticos e de simulação demonstram ganhos significativos de confiabilidade e equidade em relação a estratégias convencionais, além de fornecer diretrizes para o agrupamento de usuários e o dimensionamento de sistemas HRSMA em redes sem fio de próxima geração.

No Capítulo 16, avalia-se o desempenho de *Non-Terrestrial Networks* (NTN) baseadas em NOMA sob desvanecimento generalizado, em um cenário assistido por *Unmanned Aerial Vehicle* (UAV) com usuários terrestres. Adota-se um modelo de canal capaz de representar não linearidades, múltiplos *clusters* de propagação e correlação entre componentes, do qual se derivam expressões analíticas para a probabilidade de indisponibilidade, validadas por simulações de Monte Carlo. Os resultados mostram que o ganho de diversidade cresce com a não line-

aridade e o número de componentes Gaussianas, e que a alocação de potência é crucial sob interferência residual do *Successive Interference Cancellation* (SIC), reforçando a importância de modelos de canal realistas.

No Capítulo 17, retoma-se o mesmo cenário de NTN-NOMA assistido por UAV sob desvanecimento generalizado, avançando da caracterização para a otimização da alocação de potência. Formula-se um problema Max-Min entre as taxas dos dois usuários e resolve-se via PSO com parada antecipada, superando a limitação da alocação fixa considerada anteriormente. Os resultados mostram convergência rápida do algoritmo e equilíbrio quase ótimo entre as taxas, mas revelam um piso de desempenho em alto *Signal-to-Noise Ratio* (SNR) imposto pelas imperfeições de hardware e do SIC. Além disso, identifica-se o limiar de erro de SIC acima do qual o NOMA passa a ter desempenho inferior ao acesso ortogonal (*Orthogonal Multiple Access* (OMA)).

2 Projeto Conjunto do *Beamforming* Analógico e da BD-RIS para Transferência Eficiente de Energia Sem Fio

Victoria D. P. Souto, Osmel M. Rosabal, Onel L. A. López, Richard D. Souza, Bartolomeu F. Uchôa-Filho, Hirley Alves

2.1 Introdução

Requisitos rigorosos de confiabilidade, resiliência e eficiência energética vêm sendo estabelecidos para atender às demandas das futuras redes de IoT [6]. Nesse contexto, a 6G visa viabilizar redes IoT sustentáveis e energeticamente eficientes [6]. Entretanto, a alimentação contínua de um grande número de dispositivos ainda representa um desafio. Uma solução promissora é a WET via RF [6, 7]. A WET é considerada uma tecnologia-chave para as futuras redes IoT, pois reduz a dependência de baterias, diminui custos financeiros e ambientais, prolonga a vida útil dos dispositivos e viabiliza aplicações como operação sem bateria e dispositivos sem portas físicas [6]. Contudo, sua adoção em larga escala requer ampla cobertura de carregamento, a qual depende do projeto de *beamforming* nos nós de carregamento, usualmente denominados PBs [6, 7, 8].¹

Outra tecnologia amplamente investigada são as RIS, metasuperfícies planares compostas por elementos passivos capazes de controlar o canal sem fio por meio do ajuste de fase [9]. A integração de RIS em sistemas de comunicação sem fio pode melhorar cobertura, capacidade e eficiência energética [9]. Em sistemas WET, essa integração já demonstrou potencial para ampliar a cobertura de carregamento e o número de dispositivos IoT suportados [10, 11, 12]. No entanto, RISs convencionais controlam apenas a fase do sinal incidente, limitando seus graus de liberdade e sua capacidade de *beamforming* [9].

Recentemente, foram propostas arquiteturas denominadas BD-RIS, nas quais a matriz de espalhamento pode ser não diagonal, permitindo o controle conjunto da fase e da magnitude do sinal incidente [13, 14]. Essas arquiteturas generalizam o modelo clássico de RIS ao interconectar elementos, ou grupos de elementos, por meio de redes de impedâncias reconfiguráveis, resultando em estruturas totalmente conectadas ou conectadas por grupos [13].

Trabalhos recentes têm investigado a integração de BD-RIS em sistemas de comunicação sem fio. Por exemplo, em [13], os autores demonstram que, ao projetar adequadamente as fases e magnitudes dos elementos da BD-RIS, a potência do sinal recebido pode ser aumentada em comparação com a RIS tradicional. Em [15], os autores avaliam o problema de maximização da potência do sinal recebido em sistemas *Single-Input Single-Output* (SISO), *Multiple-Input Multiple-Output* (MIMO) e *Multiple-Input Single-Output* (MISO) auxiliados por uma BD-RIS. Seus resultados indicam que a BD-RIS pode melhorar significativamente a potência do sinal recebido em ambos os cenários quando comparada à RIS tradicional. Além disso, em [16, 17], os autores mostram que a potência do sinal recebido pode ser aumentada com a arquitetura conectada por grupos. Notavelmente, em [17], é demonstrado que a BD-RIS totalmente conectada pode alcançar desempenho quase ótimo usando apenas um bit de quantização, enquanto a RIS clássica requer quatro bits. Em [18], os autores propõem uma abordagem unificada para oti-

¹O conteúdo desta seção consiste em um resumo adaptado pela pesquisadora *Victoria D. P. Souto*, com base no artigo publicado por V. D. P. Souto, O. M. Rosabal, O. L. A. López, R. D. Souza, B. F. Uchôa-Filho and H. Alves, “Joint Analog Beamforming and Beyond-Diagonal RIS Design for Efficient Wireless Energy Transfer,” *2025 Symposium on Internet of Things (SIoT)*, Campinas, Brazil, 2025, pp. 1-4, doi: 10.1109/SIoT68426.2025.11368806.

mizar a configuração da BD-RIS considerando diversas restrições não convexas. Eles abordam os problemas de minimização de potência e maximização da eficiência energética, destacando a superioridade da BD-RIS em comparação à RIS clássica em termos de consumo de potência e eficiência energética.

A maioria dos estudos existentes emprega algoritmos iterativos, como otimização alternada e *Successive Convex Approximation* (SCA), cuja convergência depende de inicialização adequada e da solução de múltiplos subproblemas [18, 19]. Apesar dos avanços obtidos em diferentes contextos, até onde é de conhecimento dos autores, a integração entre BD-RIS e sistemas WET ainda não foi investigada na literatura. Nesse sentido, este trabalho avalia o desempenho de sistemas WET auxiliados por BD-RIS. Para isso, propõe-se uma solução de *beamforming* baseada em PSO, visando maximizar a energia coletada por dispositivos IoT sob as restrições impostas pela BD-RIS. A abordagem é analisada considerando arquiteturas de conexão única, totalmente conectada e conectada por grupos. Os resultados mostram que, mesmo com a relação não linear entre a potência de RF recebida e a energia coletada, a integração de BD-RIS em sistemas WET pode aumentar a energia coletada em pelo menos 35% com a arquitetura totalmente conectada e em até 30% com a arquitetura conectada por grupos, em comparação à RIS clássica. Além disso, a arquitetura conectada por grupos proporciona ganho de pelo menos 80% em relação ao cenário “Sem RIS”, mesmo com número reduzido de elementos. Esses resultados evidenciam o potencial da integração entre WET e BD-RIS para redes IoT mais eficientes e sustentáveis.

2.2 Modelo do Sistema

Conforme ilustrado na Figura 1, consideramos um sistema WET composto por um PB com N antenas dispostas em um *Uniform Linear Array* (ULA), auxiliado por uma BD-RIS com M elementos refletorres, responsável por alimentar um dispositivo IoT de antena única. Assume-se mobilidade limitada, de modo que os canais variam lentamente. Neste trabalho, adota-se *beamforming* analógico no PB devido à sua menor complexidade de hardware e menor consumo energético em comparação às arquiteturas digital e híbrida [20].

Logo, a potência de RF recebida pelo dispositivo IoT é expressa por

$$P = \left| \left(\mathbf{h}_R \mathbf{\Phi} \mathbf{H}_G + \mathbf{h}_D \right) \mathbf{w} \right|^2. \quad (1)$$

em que $\mathbf{h}_R \in \mathbb{C}^{1 \times M}$ representa o canal entre a BD-RIS e o dispositivo, $\mathbf{\Phi} \in \mathbb{C}^{M \times M}$ é a matriz de coeficientes de reflexão da BD-RIS, $\mathbf{H}_G \in \mathbb{C}^{M \times N}$ denota o canal entre o PB e a BD-RIS, $\mathbf{h}_D \in \mathbb{C}^{1 \times N}$ representa o canal direto entre o PB e o dispositivo e $\mathbf{w} \in \mathbb{C}^{N \times 1}$ é o vetor de *beamforming* no PB.

A energia coletada pelo dispositivo é modelada como [21]

$$E = T_c \left[\frac{\mu(1 - \epsilon\Omega)}{\epsilon(1 - \Omega)} \right], \quad (2)$$

em que T_c é o tempo de coerência do canal, $\Omega = \frac{1}{1 + e^{ab}}$ garante a condição de entrada nula/saída nula, $\epsilon = [1 + e^{-a(P-b)}]$, a e b são parâmetros do circuito [22], e μ é a potência máxima coletada no ponto de saturação do circuito. Assim, (2) descreve uma relação não linear entre a potência de RF recebida e a energia coletada, compatível com sistemas WET práticos, como os módulos *Powercast Powerharvester*® [23].

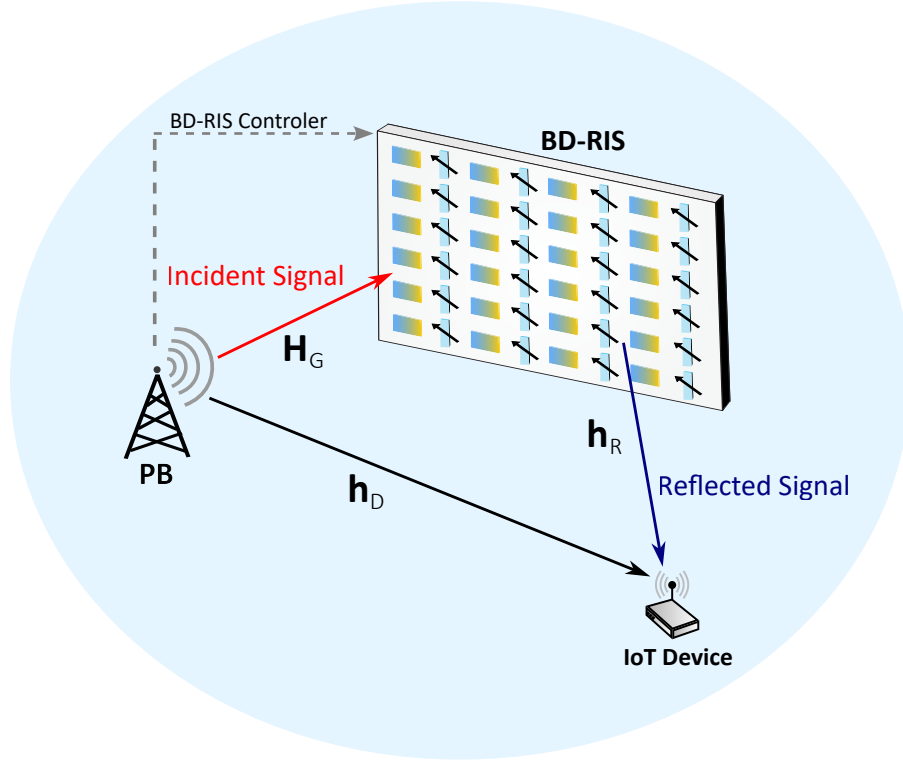


Figura 1: Modelo do sistema: um PB com N antenas, auxiliado por uma BD-RIS com M elementos, alimenta um dispositivo de antena única.

2.2.1 Arquiteturas de BD-RIS

A definição de Φ depende da arquitetura da BD-RIS. De modo geral, a BD-RIS é modelada como uma rede reconfigurável de parâmetros de espalhamento com M portas, usualmente implementada por meio de uma rede de impedâncias [13, 14]. A topologia de interconexão dessas portas determina as restrições impostas a Φ . Neste trabalho, consideram-se três arquiteturas.

- **BD-RIS de conexão única (RIS clássica):** corresponde à RIS convencional, na qual cada porta é independente e conectada ao terra por uma impedância reconfigurável [13]. Nesse caso, Φ é diagonal:

$$\Phi = \text{diag}(e^{j\varphi_1}, \dots, e^{j\varphi_M}), \quad (3)$$

$$|e^{j\varphi_m}| = 1 \quad \forall m \in \{1, \dots, M\},$$

em que $\varphi = [\varphi_1, \dots, \varphi_M]$ é o vetor de fases, com $\varphi_m \in [0, 2\pi)$. Nessa arquitetura, os elementos não diagonais de Φ são nulos, e os elementos diagonais possuem módulo unitário para satisfazer a conservação de energia [16].

- **BD-RIS totalmente conectada (fully-connected):** nessa arquitetura, todas as M portas são interconectadas por impedâncias ajustáveis [13]. Assim, Φ satisfaz

$$\Phi = \Phi^T, \quad \Phi\Phi^H = \mathbf{I}_M. \quad (4)$$

Como Φ deixa de ser diagonal, aumentam-se os graus de liberdade disponíveis para o projeto de *beamforming*. A BD-RIS de conexão única (*single connected*) pode ser vista como um caso particular da arquitetura totalmente conectada.

- **BD-RIS conectada por grupos (*group connected*):** nessa arquitetura, as M portas são divididas em N_G grupos, cada um com $M_G = M/N_G$ portas. As portas de um mesmo grupo são interconectadas, enquanto grupos distintos permanecem independentes [13]. Assim,

$$\begin{aligned} \Phi &= \text{blockdiag}(\Phi_1, \dots, \Phi_{N_G}), \\ \Phi_g &= \Phi_g^T, \quad \Phi_g \Phi_g^H = \mathbf{I}_{M_G}, \quad \forall g \in \{1, \dots, N_G\}. \end{aligned} \quad (5)$$

Essa arquitetura reduz-se à BD-RIS de conexão única quando $N_G = M$ e à BD-RIS totalmente conectada quando $N_G = 1$ [13]. Neste trabalho, utiliza-se (5) como modelo geral para Φ , variando-se N_G para contemplar as diferentes arquiteturas.

Por fim, este artigo tem como objetivo maximizar a energia coletada pelo dispositivo IoT por meio da otimização conjunta do vetor de *beamforming* no PB, $\mathbf{w}$, e da matriz de coeficientes de reflexão da BD-RIS, Φ . Essa otimização considera simultaneamente os requisitos do *beamforming* de energia analógico e as restrições impostas pela arquitetura da BD-RIS. Para resolver o problema de otimização proposto, propõe-se uma solução baseada em PSO, uma técnica de otimização estocástica reconhecida por sua baixa complexidade computacional e elevada efetividade na solução de problemas não convexos, como o considerado neste trabalho [24, 25]. Os resultados obtidos a partir da solução proposta são apresentados a seguir.

2.3 Resultados de Simulação

A avaliação da abordagem proposta considera o cenário da Figura 1. Nesta seção, analisa-se o desempenho do sistema WET auxiliado por BD-RIS de conexão única (*Classical RIS*), conectada por grupos (*group connected*) e totalmente conectada (*fully connected*). Como referência, considera-se também o cenário sem RIS, no qual o *beamforming* no PB é ótimo. Salvo indicação em contrário, os parâmetros de simulação são $M = 30$, $N = 4$, $T_c = 2$ s, $f_c = 915$ MHz, $\mu = 10.73$ mW, $a = 5.365$, $b = 0.2308$, $P_T = 3$ W, $[\kappa_R, \kappa_D, \kappa_G] = [1.0, 1.0, 1.0]$, $[\alpha_R, \alpha_D, \alpha_G] = [2.0, 3.5, 2.0]$, $[x_{PB}, y_{PB}] = [0, 0]$ m, $[x_{RIS}, y_{RIS}] = [10, 3]$ m e $[x_D, y_D] = [12, 0]$ m. Todas as curvas correspondem à média de 10^3 realizações independentes de canal.

A Figura 2 apresenta a energia coletada para diferentes valores de M . Observa-se que, mesmo para valores reduzidos de M , as arquiteturas BD-RIS totalmente conectada e conectada por grupos superam a RIS clássica. A BD-RIS totalmente conectada proporciona ganhos mínimos de 14%, 21%, 30%, 32% e 36% para $M \in \{10, 20, 30, 40, 50\}$, respectivamente. Já a arquitetura conectada por grupos, com $N_G = 5$, alcança ganhos mínimos de 5%, 16%, 23%, 28% e 32% para os mesmos valores de M . Além disso, mesmo com uma BD-RIS conectada por grupos com $N_G = 5$, obtém-se um aumento de pelo menos 80% na energia coletada em relação ao cenário sem RIS. Esses resultados evidenciam a relevância de arquiteturas avançadas de RIS para o aumento da eficiência de sistemas WET nas futuras redes IoT.

A Figura 3 mostra a energia coletada em função do número de antenas no PB. A BD-RIS totalmente conectada apresenta o melhor desempenho em todos os casos avaliados. No entanto, à medida que N aumenta, a diferença relativa em relação à RIS clássica diminui, devido à maior contribuição do enlace direto PB-dispositivo. Em particular, a arquitetura totalmente conectada proporciona ganhos mínimos de 31%, 30%, 26% e 25% para $N \in \{2, 4, 6, 8\}$, respectivamente. Para a BD-RIS conectada por grupos com $N_G = 5$, os ganhos mínimos são de 27%, 23%, 22% e 19%. Além disso, em comparação ao cenário sem RIS, o uso de BD-RIS pode

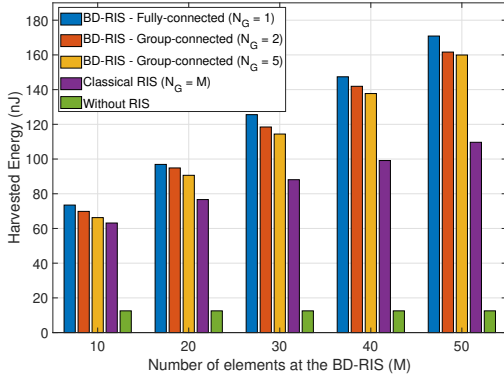


Figura 2: Energia coletada em função do número de elementos M .

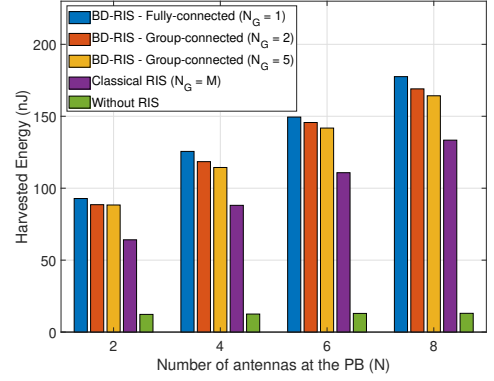


Figura 3: Energia coletada em função do número de antenas N .

aumentar a energia coletada em pelo menos 90%, mesmo para valores elevados de N . Portanto, os resultados confirmam a superioridade das arquiteturas BD-RIS em sistemas WET.

A Figura 4 apresenta a energia coletada para diferentes valores do fator Rician no enlace direto PB–dispositivo (κ_D). Observa-se que a energia coletada permanece aproximadamente estável com a variação de κ_D . Como esperado, a BD-RIS totalmente conectada supera a RIS clássica e o cenário sem RIS. Essa superioridade também se mantém na presença de um enlace direto com $\kappa_D = 0.5$ e $\alpha_D = 2.0$ (resultado omitido por limitação de espaço). Tal comportamento decorre da capacidade da PSO de otimizar o *beamforming* de transmissão considerando conjuntamente os enlaces direto PB–dispositivo e refletido PB–BD-RIS–dispositivo, aumentando o desempenho do sistema WET.

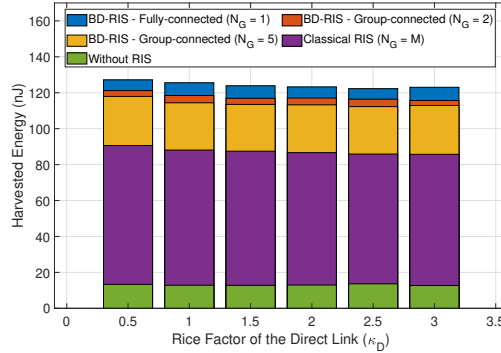


Figura 4: Energia coletada vs. fator de Rician κ_D .

2.4 Conclusões

Propusemos um método baseado em PSO para maximizar a energia coletada por meio da otimização conjunta do *beamforming* no PB e da matriz de coeficientes de reflexão na BD-RIS, considerando arquiteturas de conexão única, totalmente conectada e conectada por grupos. A solução proposta comprova que a integração de BD-RIS em sistemas WET pode melhorar significativamente a energia coletada, mesmo com a arquitetura conectada por grupos, que apresenta menor complexidade de *hardware*. Em trabalhos futuros, pretendemos investigar a integração de WET e BD-RIS em cenários multiusuário.

3 Análise de Pré-Distorção Digital em Sistemas FTN-GFDMA Integrados ao A-RoF

Mariana B. de Mello, Luiz A. M. Pereira, Luciano L. Mendes
Inatel

3.1 Introdução

Neste trabalho foi analisada a integração da técnica *Faster-than-Nyquist* (FTN)-*Generalized Frequency Division Multiple Access* (GFDMA) em sistemas *Analog Radio over Fiber* (A-RoF), com foco em aplicações *Enhanced Remote Areas Communication* (eRAC). A utilização do FTN-GFDMA em conjunto com arquiteturas A-RoF permite aumentar a eficiência espectral e a flexibilidade do sistema, possibilitando maior capacidade de transmissão mesmo em cenários com limitação de banda disponível [26, 27].

O crescimento das futuras redes 6G demanda soluções capazes de oferecer altas taxas de transmissão, baixa latência e ampla cobertura, incluindo regiões rurais e remotas. Nesse contexto, arquiteturas A-RoF associadas a redes centralizadas surgem como alternativas promissoras para reduzir custos de infraestrutura e ampliar a cobertura de sistemas sem fio. Além disso, o uso do FTN-GFDMA possibilita melhor aproveitamento espectral por meio da sobreposição controlada dos símbolos no domínio do tempo e da frequência, aumentando a densidade de símbolos transmitidos sem necessidade de expansão significativa da largura de banda disponível [28, 29].

Outra característica importante do FTN-GFDMA é sua capacidade de reduzir *Out-of-Band Emissions* (OOBE) por meio da utilização de filtros protótipos. Essa propriedade torna a técnica especialmente interessante para aplicações em faixas espectrais limitadas, como canais TVWS, favorecendo a coexistência com outros sistemas de comunicação e aumentando a flexibilidade na alocação de recursos [30].

Entretanto, a integração entre FTN-GFDMA e A-RoF introduz desafios importantes relacionados às não linearidades do sistema óptico, principalmente devido ao comportamento não linear do *Mach-Zehnder Modulator* (MZM). Em condições práticas de operação, especialmente em enlaces de longa distância e cenários com alocação dinâmica de potência, o MZM pode operar em regiões não lineares, causando degradação do sinal transmitido, aumento das emissões espectrais indesejadas e elevação da BER.

Além disso, o próprio FTN-GFDMA apresenta desafios relacionados ao aumento da *Intersymbol Interference* (ISI) e *Intercarrier Interference* (ICI), devido à violação intencional do critério de Nyquist [29][31]. Essas interferências, somadas aos efeitos não lineares do enlace óptico, tornam necessária a utilização de técnicas de compensação e linearização para garantir desempenho adequado do sistema.

Para mitigar esses efeitos, foi empregada uma técnica de *Digital Pre-Distortion* (DPD) baseada em aprendizado de máquina utilizando redes neurais artificiais do tipo *Multi-Layer Perceptron* (MLP). O objetivo da DPD é aprender uma função inversa da não linearidade introduzida pelo transmissor óptico, reduzindo as distorções antes da transmissão do sinal pelo enlace A-RoF [32, 33, 34]. Dessa forma, busca-se melhorar o desempenho do sistema em termos de BER, aproximando os resultados obtidos do comportamento teórico esperado para sistemas lineares.

3.2 Digital Pre-Distortion Baseado em Redes Neurais

A presença de não linearidades na cadeia de transmissão óptica representa um dos principais desafios para a integração do FTM-GFDMA em sistemas A-RoF. Em particular, o MZM apresenta resposta não linear quando operando fora de sua região linear de transferência, condição que pode ocorrer em cenários com maiores níveis de potência elétrica de excitação ou elevada flutuação instantânea de amplitude do sinal transmitido. Como consequência, surgem distorções que degradam o desempenho do sistema em termos de BER e comprometem a integridade do sinal transmitido.

Esse problema torna-se ainda mais crítico em sinais FTM-GFDMA, uma vez que a combinação entre múltiplas subportadoras, sobreposição temporal e filtragem protótipo pode aumentar a PAPR do sinal transmitido, tornando-o mais suscetível às não linearidades do MZM. Além disso, os efeitos de ISI e ICI inerentes ao FTM-GFDMA, decorrentes da violação controlada do critério de Nyquist, aumentam a sensibilidade do sistema às degradações introduzidas pelo enlace óptico. Dessa forma, técnicas de linearização tornam-se fundamentais para viabilizar a transmissão de sinais FTM-GFDMA em arquiteturas A-RoF mantendo desempenho adequado de detecção ².

Com o objetivo de mitigar essas distorções, foi empregada uma técnica de DPD baseada em aprendizado de máquina utilizando redes neurais artificiais do tipo MLP. O princípio do DPD consiste em aplicar uma pré-distorção ao sinal em banda base antes da modulação eletro-óptica, de forma que a resposta não linear introduzida pelo MZM seja previamente compensada. Para isso, a rede neural foi treinada de forma supervisionada utilizando conjuntos de dados experimentais obtidos em uma cadeia real de transmissão A-RoF. Nesse processo, as amostras distorcidas coletadas após o estágio de modulação óptica foram utilizadas como entrada da rede, enquanto as amostras originais do sinal em banda base foram utilizadas como sinais de referência. Dessa forma, a rede neural aprende uma aproximação inversa da resposta não linear da cadeia óptica, permitindo gerar um sinal pré-distorcido que reduz as distorções introduzidas pelo MZM e aproxima o comportamento do sistema ao de uma cadeia linear [32].

O bloco de DPD foi implementado utilizando uma rede neural *feedforward* do tipo MLP, treinada para mapear amostras distorcidas do sinal em banda base em suas correspondentes versões não distorcidas. A rede neural processa entradas no domínio *real-valued*, obtidas a partir da representação complexa do sinal, e gera como saída as componentes em fase e quadratura do sinal pré-distorcido.

A arquitetura adotada é composta por três camadas ocultas contendo 288, 544 e 224 neurônios, respectivamente. As funções de ativação empregadas nas camadas ocultas foram *Rectified Linear Unit* (ReLU), linear e *Scaled Exponential Linear Unit* (SELU). A camada de saída utiliza ativação linear e contém dois neurônios responsáveis pela geração das componentes real e imaginária do sinal pré-distorcido. A arquitetura da rede neural utilizada como bloco de DPD é ilustrada na Fig. 5.

Todos os parâmetros da rede foram inicializados utilizando uma distribuição Gaussiana de média zero e desvio padrão igual a 0,05, de forma a garantir convergência estável durante o treinamento.

O treinamento foi realizado de forma supervisionada utilizando conjuntos de dados experimentais em banda base obtidos de uma cadeia de transmissão A-RoF. As amostras distorcidas

²Mais detalhes sobre a modelagem matemática e a integração do sistema FTM-GFDMA com a arquitetura A-RoF podem ser encontrados na Parte I do relatório da Atividade 2.2 – Concepção da Rede de Acesso para Rede 6G, referente à Fase III.

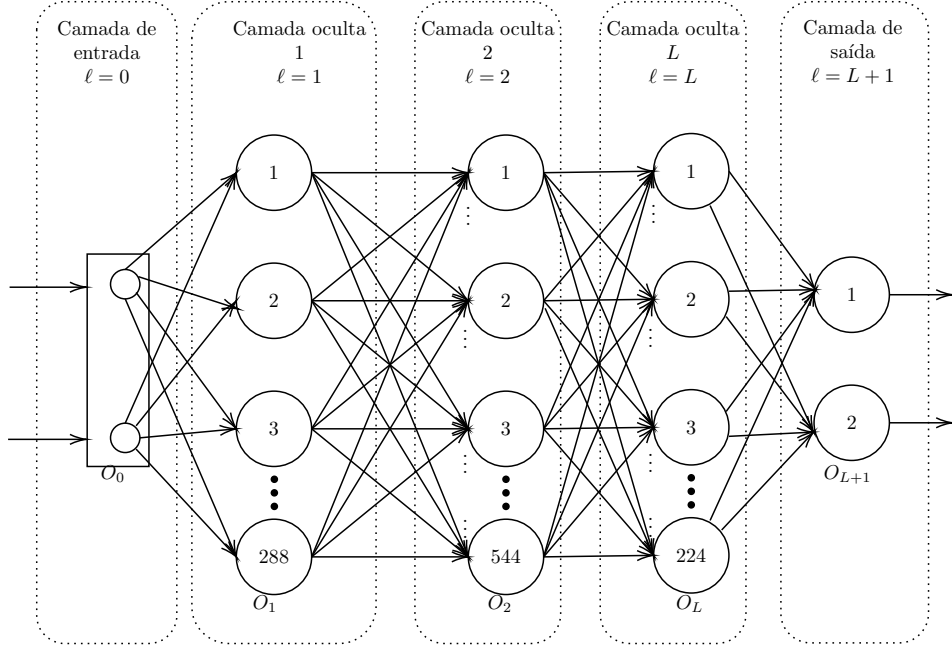


Figura 5: Arquitetura da rede neural artificial do tipo multi-layer perceptron utilizada como bloco de DPD.

coletadas após os estágios do MZM foram utilizadas como entradas da rede, enquanto as correspondentes amostras originais não distorcidas do sinal em banda base foram utilizadas como sinais de referência.

Antes do treinamento, os sinais de entrada e referência foram normalizados para apresentar mesma potência média, evitando instabilidades numéricas e viés devido ao escalonamento de amplitude. O conjunto de dados foi dividido em subconjuntos de treinamento e validação, e a rede neural foi otimizada utilizando o algoritmo Adam com taxa de aprendizado de $3,98 \times 10^{-5}$.

Como função de custo, foi utilizado o erro quadrático médio (*Mean Squared Error* (MSE)) entre as componentes I/Q preditas e os sinais de referência. O treinamento foi realizado por até 500 épocas utilizando batches contendo 4096 amostras. Além disso, foi utilizado um critério de *early stopping* para evitar *overfitting*, restaurando os parâmetros da rede correspondentes ao menor erro de treinamento obtido.

Após o treinamento, a rede neural utilizada como bloco de DPD foi inserida na cadeia de transmissão do sistema FTN-GFDMA, conforme ilustrado na Fig. 6. O sinal pré-distorcido $\mathbf{x}'$ é obtido aplicando o método *predict* da rede neural baseada em *Artificial Neural Network* (ANN) utilizando $\mathbf{x}$ como entrada.

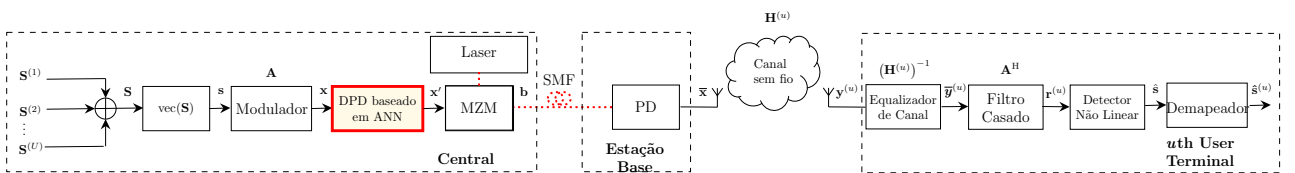


Figura 6: Diagrama de blocos do transceptor downlink FTN-GFDMA integrado ao enlace A-RoF, incluindo o bloco de DPD baseado em ANN na central office/transmissor.

Uma vantagem importante dessa abordagem é que o bloco de DPD é implementado na

Central Office (CO), preservando a baixa complexidade das unidades remotas do sistema A-RoF. Essa característica é particularmente relevante em aplicações eRAC, nas quais a redução de custos de implantação e manutenção da infraestrutura representa um requisito importante.

Além disso, a utilização de redes neurais permite modelar a resposta não linear do transmissor óptico sem necessidade de uma modelagem analítica explícita do sistema. Dessa forma, o DPD baseado em aprendizado de máquina apresenta-se como uma solução promissora para reduzir as distorções não lineares em sistemas FTN-GFDMA integrados a enlaces A-RoF, contribuindo para melhorar o desempenho do sistema em termos de BER.

3.3 Resultados e Discussões

Esta seção apresenta os resultados relacionados à avaliação da técnica de DPD aplicada ao sistema FTN-GFDMA integrado ao enlace A-RoF. Como as análises de BER do sistema FTN-GFDMA e de PAPR já foram discutidas anteriormente, esta seção concentra-se nos impactos das não linearidades introduzidas pela cadeia óptica e na capacidade do DPD baseado em ANN em compensar essas distorções.

Os resultados foram obtidos considerando modulação *Binary Phase Shift Keying* (BPSK) e filtro *Radio Resource Control* (RRC). Os parâmetros do sistema utilizados nas simulações são apresentados na Tabela 1. Além disso, a Tabela 2 apresenta o modelo de canal adotado para a avaliação do desempenho do sistema.

Tabela 1: Parâmetros do sistema FTN-GFDM utilizados na avaliação do DPD.

Parâmetro	Valor
Modulação	BPSK
Filtro protótipo	RRC
Fator de roll-off (α)	0
Número de períodos do filtro protótipo ($\mathcal{P}$)	4
Número de amostras por período do filtro ($\mathcal{S}$)	5
Espaçamento temporal entre subsímbolos (Δ_K)	4
Espaçamento em frequência entre subportadoras (Δ_M)	4
Fator de sobreposição temporal (τ)	0,8
Fator de sobreposição espectral (ϕ)	1
Subportadoras por bloco (K)	5
Subsímbolos por bloco (M)	5

Tabela 2: Modelo de canal utilizado na avaliação do DPD.

Canal	Resposta ao impulso
AWGN	$h = 1$

A avaliação do desempenho do DPD segue o diagrama de blocos apresentado na Fig. 6. Considerando a hipótese de canal *Additive White Gaussian Noise* (AWGN), sem desvanecimento multipercurso, a matriz de canal $\mathbf{H}$ foi definida como uma matriz identidade.

A Fig. 7 apresenta o desempenho de BER em função de E_b/N_0 para o sinal FTN-GFDM-I transmitido através do enlace A-RoF, considerando os cenários com e sem a utilização do bloco de DPD. Além disso, a curva teórica de BPSK em canal AWGN é apresentada como referência.

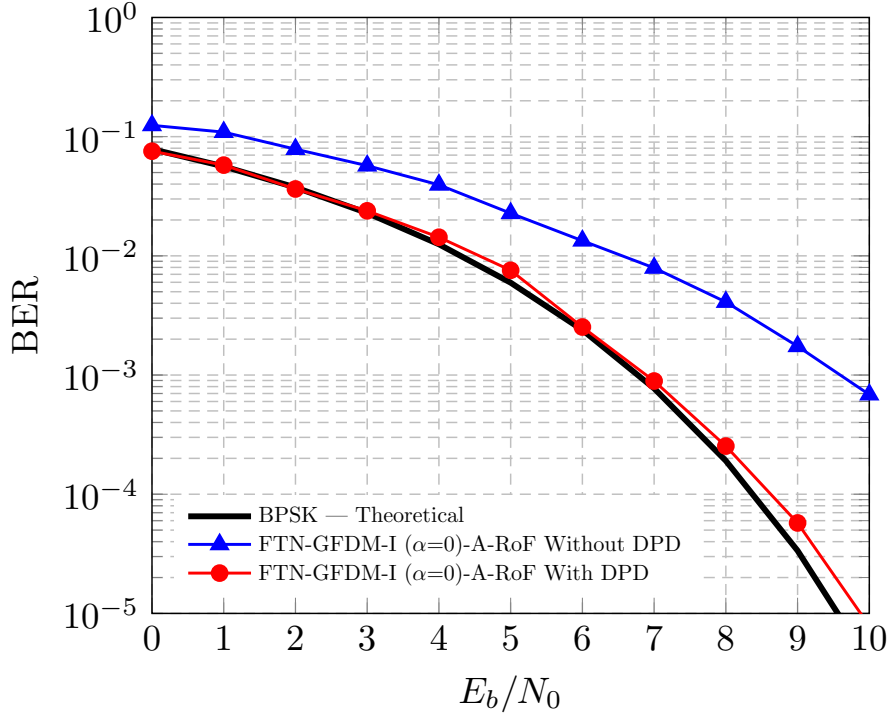


Figura 7: Desempenho de BER do sistema FTN-*Generalized Frequency Division Multiplexing* (GFDM) transmitido sobre um enlace A-RoF, considerando cenários com e sem DPD.

Observa-se que o sistema operando sem DPD apresenta degradação significativa de desempenho em toda a faixa de relação sinal-ruído analisada. Esse comportamento ocorre devido às distorções não lineares introduzidas pelo MZM, que afetam diretamente a integridade do sinal transmitido e aumentam a probabilidade de erro na detecção dos símbolos.

Quando o bloco de DPD baseado em ANN é aplicado, observa-se uma melhoria expressiva no desempenho de BER. A curva correspondente ao sistema com DPD aproxima-se significativamente da curva teórica de BPSK, principalmente na região de baixos e moderados valores de E_b/N_0 , indicando que a rede neural foi capaz de compensar grande parte das distorções não lineares introduzidas pela cadeia A-RoF.

A redução observada na BER demonstra que o modelo baseado em ANN consegue aprender uma aproximação inversa da resposta não linear do transmissor óptico, mitigando as distorções introduzidas antes da modulação eletro-óptica. Dessa forma, o DPD contribui para aproximar o comportamento do sistema ao de uma cadeia linear ideal.

Em valores elevados de E_b/N_0 , ainda é possível observar uma pequena diferença entre a curva obtida com DPD e o limite teórico. Esse comportamento pode estar associado a efeitos não modelados pela rede neural, como componentes não lineares de ordem superior e distorções residuais presentes na cadeia óptica.

De forma geral, os resultados demonstram que a utilização de DPD baseado em aprendizado de máquina representa uma solução eficiente para reduzir os efeitos das não linearidades em sistemas FTN-GFDMA integrados a enlaces A-RoF, permitindo explorar os ganhos de eficiência espectral da técnica sem degradação significativa do desempenho em termos de BER.

3.4 Conclusão

Este trabalho investigou a utilização de uma técnica de DPD baseada em aprendizado de máquina para compensar as não linearidades introduzidas pela cadeia de transmissão A-RoF em sistemas FTN-GFDMA. Considerando o alto nível de PAPR e a maior sensibilidade do FTN-GFDMA às distorções não lineares do MZM, a aplicação de técnicas de linearização torna-se fundamental para preservar o desempenho do sistema.

Os resultados obtidos mostraram que o bloco de DPD baseado em ANN foi capaz de reduzir significativamente os efeitos das não linearidades introduzidas pelo transmissor óptico. A compensação realizada pela rede neural permitiu aproximar o desempenho de BER do sistema da curva teórica de referência em canal AWGN, demonstrando a capacidade da técnica em aprender uma aproximação inversa da resposta não linear da cadeia A-RoF.

Além da melhoria de desempenho, a implementação do DPD na CO preserva a baixa complexidade das unidades remotas, característica importante para aplicações eRAC. Dessa forma, a utilização de técnicas baseadas em aprendizado de máquina surge como uma alternativa promissora para viabilizar a integração entre FTN-GFDMA e sistemas A-RoF, permitindo explorar os ganhos de eficiência espectral da forma de onda sem degradação significativa causada pelas não linearidades ópticas.

Como trabalhos futuros, pretende-se investigar técnicas conjuntas de redução de PAPR e linearização, além da avaliação do sistema em cenários com enlaces ópticos de maior distância, modulações de ordem superior e cenários multiusuário.

4 Aprendizado por Reforço para Reduzir PAPR em Sistemas de Comunicação

Bianca S. de C. da Silva, Luiz A. M. Pereira, Luciano L. Mendes
Inatel

4.1 Introdução

A crescente demanda por aplicações de alta taxa de dados, como videoconferência, transmissões ao vivo e serviços de *streaming*, ampliou o número de usuários em sistemas modernos de comunicação sem fio [35]. Esse cenário intensifica o problema da escassez de espectro, especialmente nas faixas tradicionalmente destinadas aos serviços móveis, cuja alocação fixa e pouco eficiente contribui para a subutilização de recursos [36, 37].

Para mitigar essas limitações, diversos países têm adotado o uso de faixas subutilizadas do espectro, conhecidas como TVWS. No Brasil, a Anatel regulamenta o uso secundário de canais ociosos de TV em *Very High Frequency* (VHF) e *Ultra High Frequency* (UHF), ampliando a disponibilidade espectral, sobretudo em áreas rurais. Contudo, sistemas em TVWS enfrentam limites rígidos de potência, tipicamente até 1 W, para evitar interferência em serviços primários, o que impõe desafios adicionais ao desempenho dos transmissores [38, 39].

Entre as técnicas de modulação adequadas a esse contexto, o OFDM se destaca pela eficiência espectral e robustez ao desvanecimento seletivo, embora apresente o problema intrínseco do alto PAPR [40]. Esse efeito exige amplificadores altamente lineares, sob pena de gerar distorções, ICI e aumento das emissões fora da banda.

A redução do PAPR é essencial em diferentes cenários modernos, incluindo *Long-Term Evolution* (LTE)-*Unlicensed* (U), *New Radio* (NR)-U [41], aplicações *Industrial, Scientific and Medical* (ISM) [42], sistemas via satélite e redes de próxima geração. A literatura apresenta diversas técnicas de mitigação, como PTS, *Selected Mapping* (SLM) [43], filtragem e esquemas de codificação, embora muitas envolvam alta complexidade computacional. Paralelamente, métodos baseados em *Machine Learning* (ML) têm ganhado espaço por oferecer soluções mais flexíveis e, por vezes, mais eficientes.

Nesse contexto, técnicas de RL surgem como alternativas promissoras, especialmente por dispensarem bases de dados prévias. Entre elas, o algoritmo Q-Learning destaca-se pela simplicidade e por permitir a aprendizagem de estratégias por interação direta com o ambiente. Sua aplicação ao OFDM busca explorar a seleção inteligente de valores de subportadoras piloto para reduzir o PAPR, avaliando o impacto dessa abordagem em comparação com técnicas tradicionais e métodos supervisionados baseados em *Neural Network* (NN).

Na próxima seção, detalham-se o desenvolvimento da abordagem baseada em Q-Learning para o sistema OFDM e os resultados decorrentes dessa aplicação.

4.2 Desenvolvimento

Nesta seção são apresentados os fundamentos teóricos e metodológicos adotados para a análise e redução do PAPR em sistemas OFDM. São descritos o funcionamento do sistema de transmissão, o conceito de PAPR, bem como os princípios do algoritmo Q-Learning.

4.2.1 Fundamentos do Sistema OFDM

O OFDM é uma técnica amplamente utilizada em sistemas de comunicação devido à sua estrutura multicarrier e à simplicidade de processamento no domínio da frequência. Os bits de entrada, $\mathbf{b} = [b_0, b_2, \dots, b_{\mu K-1}]$, onde K é o número de subportadoras, são mapeados em símbolos complexos, $\mathbf{s} \in \mathbb{C}^{K \times 1}$, e convertidos para o domínio do tempo por meio da *Inverse Fast Fourier Transform* (IFFT):

$$\mathbf{x} = \mathbf{F}_K^H \mathbf{s}, \quad (6)$$

onde $\mathbf{F}_K$ representa a matriz de Fourier normalizada. Em seguida, insere-se o *Cyclic Prefix* (CP), composto pelos últimos K_{CP} elementos de $\mathbf{x}$. No receptor, remove-se o CP, aplica-se a *Fast Fourier Transform* (FFT) e utiliza-se uma equalização simples no domínio da frequência, por exemplo via *zero-forcing*.

Apesar de suas vantagens e larga adoção, o OFDM apresenta como uma de suas principais limitações o elevado valor de PAPR, que impacta diretamente o desempenho dos amplificadores de potência [44].

4.2.2 Caracterização do PAPR

O PAPR de um símbolo OFDM é definido como:

$$\text{PAPR} = \frac{\max\{|x[n]|^2\}}{\mathbb{E}\{|x[n]|^2\}}, \quad (7)$$

onde $\mathbb{E}\{\cdot\}$ representa o valor esperado. Esse problema ocorre porque o sinal OFDM é resultado da soma simultânea de múltiplas subportadoras, o que aumenta a probabilidade de que suas amplitudes se combinem de forma construtiva, gerando picos muito superiores ao valor médio.

Esses picos podem saturar o amplificador de potência, causando distorções, ICI e emissões fora da banda, tornando fundamental a mitigação do PAPR.

Diante desse cenário, o próximo tópico apresenta a aplicação do algoritmo mais simples de RL na mitigação desse problema.

4.2.3 Fundamentos do Q-Learning

Como forma de verificar a viabilidade do uso de RL para mitigação do PAPR em sistemas OFDM, uma abordagem ainda pouco explorada na literatura, optou-se pela adoção do algoritmo mais simples dessa classe: o Q-Learning [45, 46]. Essa escolha permite avaliar, de maneira direta, se mesmo um método elementar é capaz de aprender políticas eficazes para redução do problema.

O Q-Learning é um algoritmo baseado em tabela que estima a função de valor-ação $Q(\varsigma, a)$, a qual representa o retorno esperado ao executar a ação a no estado ς . A atualização dessa função segue a equação de Bellman [45]:

$$Q(\varsigma_t, a_t) \leftarrow Q(\varsigma_t, a_t) + \alpha \left[r_{t+1} + \gamma \max_{a'} Q(\varsigma_{t+1}, a') - Q(\varsigma_t, a_t) \right], \quad (8)$$

em que α é a taxa de aprendizado e γ o fator de desconto.

A política utilizada é do tipo ε -greedy, que equilibra *exploitation* e *exploration* ao permitir que o agente escolha ações aleatórias com probabilidade ε , e ações que maximizam Q com probabilidade $1 - \varepsilon$.

4.2.4 Aplicação do Q-Learning à Redução de PAPR

Na aplicação ao sistema OFDM, os estados do agente são definidos por faixas discretizadas de valores de PAPR obtidos a partir de símbolos OFDM sem subportadoras piloto. O intervalo típico de variação observado, entre 4 e 11 dB, é dividido em diversos estados, permitindo reduzir o espaço de busca e tornar o aprendizado viável.

As ações disponíveis ao agente correspondem às possíveis combinações de valores dos K_p pilotos, compatíveis com o espaço amostral da constelação em uso. Para cada tentativa, o agente seleciona valores de pilotos, gera o símbolo OFDM e calcula o PAPR. Caso o valor seja inferior ao limiar predefinido, o processo é finalizado. Caso contrário, novas ações são tomadas até que o critério seja satisfeito.

A Fig. 8 ilustra a integração do bloco de RL ao transmissor OFDM, destacando a busca contínua por combinações que reduzam o PAPR antes da transmissão final do símbolo.

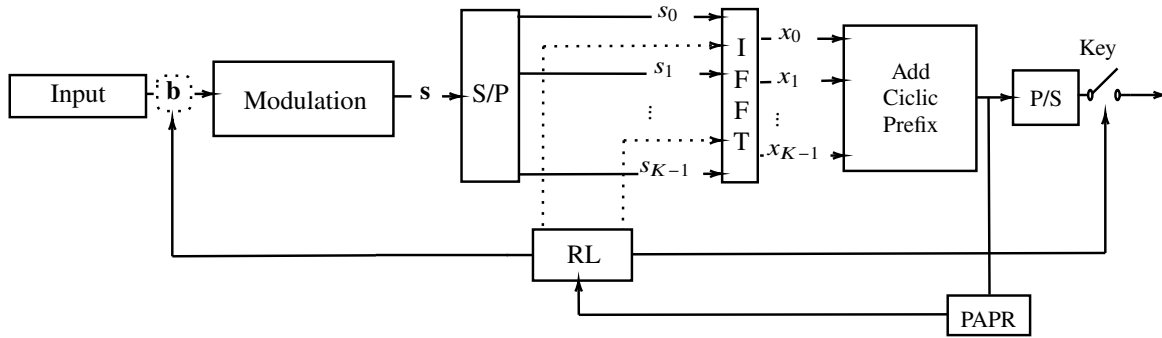


Figura 8: Transmissor OFDM com redução de PAPR baseada em RL [1].

4.3 Resultados Numéricos

Esta seção apresenta os resultados obtidos para a redução do PAPR e para o desempenho de BER dos métodos avaliados: Q-Learning, rede neural supervisionada, método de referência e PTS. As simulações consideram uma modulação 4-*Quadrature Amplitude Modulation* (QAM), $K = 64$ subportadoras, sendo $K_p = 2$ pilotos e 62 subportadoras úteis.

4.3.1 Análise de PAPR

A Fig. 9 apresenta as curvas de *Complementary Cumulative Distribution Function* (CCDF) do PAPR para os métodos avaliados. Observa-se que o Q-Learning, mesmo sendo o algoritmo de RL mais simples, apresenta desempenho superior ao PTS, reduzindo o PAPR ao longo de toda a distribuição. Esse resultado evidencia o potencial do RL em problemas clássicos de otimização em sistemas de comunicação.

A rede neural supervisionada, treinada a partir do conjunto de referência, alcança desempenho ligeiramente superior ao Q-Learning, resultado esperado devido à sua capacidade de generalização e ao uso de um dataset construído com combinações ótimas de pilotos.

4.3.2 Análise de BER

A avaliação de BER sob canal AWGN é apresentada na Fig. 10. Os resultados mostram que os métodos baseados em inserção de pilotos, Q-Learning, rede neural e referência, apresentam

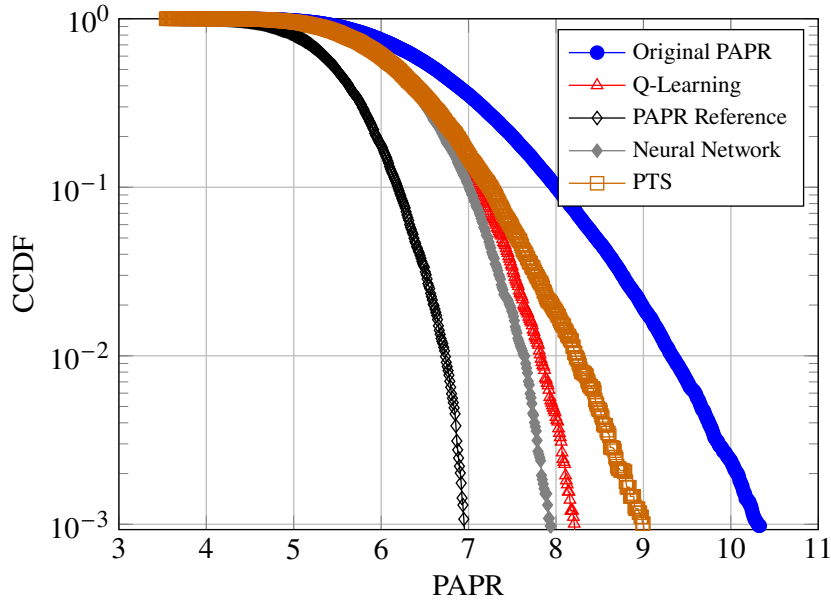


Figura 9: Curvas CCDF do PAPR para Q-Learning, rede neural, referência e PTS [1].

uma pequena degradação em relação ao BER teórico e à curva obtida para o sinal original sem pilotos.

Tal degradação ocorre devido à redução da eficiência espectral, uma vez que parte da potência total é alocada aos pilotos, reduzindo levemente a energia disponível para os símbolos úteis. Apesar disso, o impacto permanece controlado e coerente com o efeito esperado da diminuição de eficiência $\epsilon = (K - K_p)/K$.

O método PTS, embora apresente bom desempenho em BER, demonstra um comportamento de piso de erro (*error floor*) em baixas probabilidades de erro, indicando limitação estrutural na sua capacidade de melhoria.

4.4 Conclusão

Os resultados mostram que o uso de Q-Learning para redução de PAPR em sistemas OFDM é uma alternativa eficiente e de baixa complexidade frente às técnicas tradicionais. Mesmo sendo um método simples dentro de RL, o algoritmo foi capaz de aprender combinações adequadas de pilotos e reduzir o PAPR sem necessidade de bases de dados ou modelos supervisionados.

A comparação com a rede neural supervisionada confirma seu melhor desempenho, mas destaca a vantagem do Q-Learning em termos de independência de dataset e menor custo computacional. Já em relação ao método PTS, os resultados indicam superioridade do Q-Learning, reforçando o potencial das abordagens de RL para esse tipo de otimização.

A análise de BER aponta uma pequena degradação nos métodos baseados em pilotos, prevista pela redução de eficiência espectral devido à inserção dos K_p pilotos, mas sem impacto significativo no desempenho geral.

Conclui-se, portanto, que o Q-Learning é uma solução promissora para mitigação de PAPR, especialmente em cenários que demandam simplicidade e baixo custo computacional. Além disso, indica a viabilidade de investigações futuras com técnicas mais avançadas, como *Deep Q-Network* (DQN), capazes de ampliar a capacidade de generalização e explorar espaços contínuos de estados.

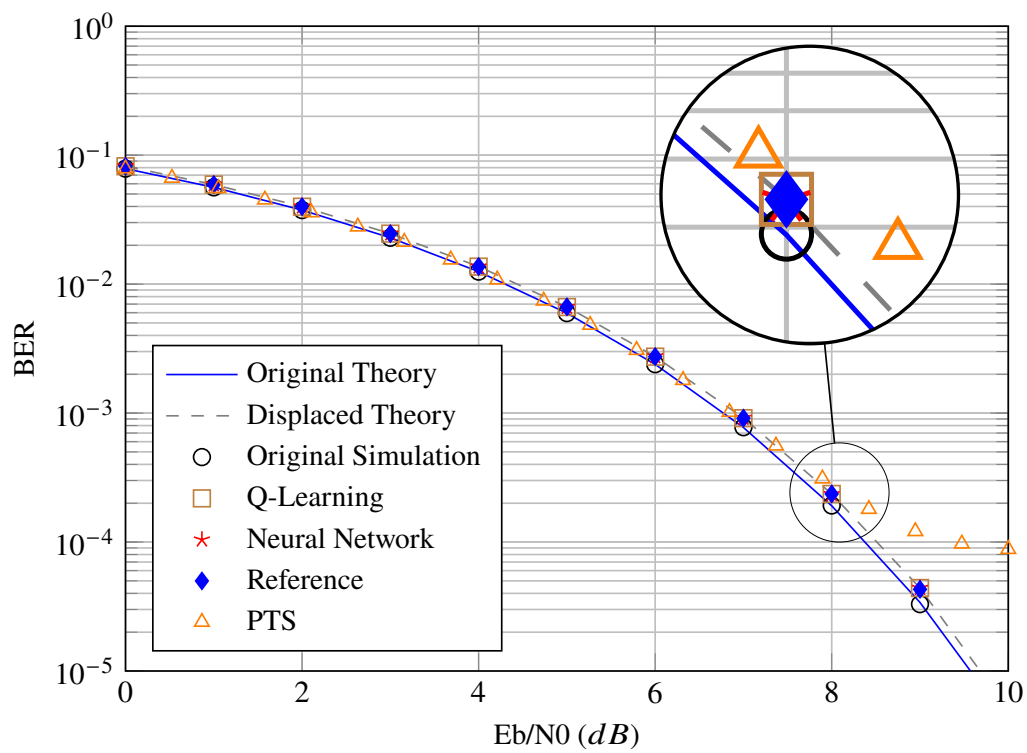


Figura 10: Desempenho de BER sob canal AWGN para os métodos analisados [1].

5 Codificação Multiterminal para Fontes Gaussianas com Ruído Não-Estacionário

Eduardo N. Velloso, José Cândido S. Santos Filho
UNICAMP

5.1 Introdução

Uma das principais ideias por trás das tecnologias habilitadoras das aplicações para 6G é a comunicação orientada a tarefas, que consiste na identificação e transmissão dos dados mais relevantes para cada tarefa em questão. A vantagem prática desse modelo reside no aumento da capacidade efetiva que ele proporciona, aproveitando qualquer contexto comum como informação adicional para melhorar a proporção de semântica (i.e., informação útil) transmitida por largura de banda disponível [47]. Para modelar a compressão orientada a tarefas, pode-se considerar que uma observação de fonte X seja correlacionada a um estado semântico intrínseco oculto S , que representa todos os aspectos da fonte relevantes para a tarefa específica, conforme ilustrado na Figura 11. Substituir o critério de fidelidade por um equivalente semântico enquanto o codificador só pode observar a fonte ruidosa transforma o problema original em um problema de codificação de fonte indireta (ou remota) com perdas.³

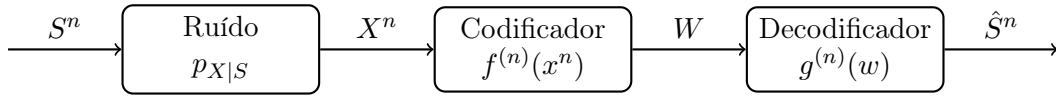


Figura 11: Compressão orientada a tarefas modelada como codificação de fonte indireta com perdas.

A codificação de fonte multiterminal, também conhecida como o problema do CEO, é uma variante da codificação de fonte indireta que a inclui como um caso especial [48]. O problema do CEO Gaussiano quadrático tem servido, já há algum tempo, como uma referência canônica para a teoria de informação multiterminal. Em sua formulação clássica, vários sensores sujeitos a ruído observam uma fonte Gaussiana oculta e transmitem descrições comprimidas a um decodificador central, o qual reconstrói a fonte sob uma restrição de erro médio quadrático. Para o caso escalar sem memória, toda a região de taxa foi caracterizada por meio de um limitante estrito de Berger-Tung [49, 50]. Esta linha de pesquisa estabeleceu tanto a taxa-soma ótima quanto uma estratégia de codificação clara baseada em binagem Slepian-Wolf [51, 52].

No caso de fontes vetoriais, ou de fontes escalares com memória, a situação é consideravelmente mais complexa. Quando a fonte e os ruídos da observação são simultaneamente diagonalizáveis, o problema se decompõe em canais escalares paralelos e se reduz a um conjunto de otimizações independentes de uma única variável. No entanto, no caso não-comutativo genérico, essa diagonalização falha, e a região de taxa exata só é conhecida em casos especiais.

Uma linha de pesquisa paralela desenvolveu métodos de tempo-frequência e de “espaço de fase” para processos Gaussianos não-estacionários. Com base nos resultados sobre a distribuição de autovalores pelo teorema de Szegő [53] e na teoria dos processos localmente estacionários

³O conteúdo desta seção consiste em um resumo adaptado pelo pesquisador *Eduardo N. Velloso*, com base no artigo publicado no IEEE International Symposium on Information Theory (ISIT 2026), intitulado “Time-Frequency Waterfilling for the Nonstationary Gaussian CEO Problem”.

[54], demonstrou-se que a função taxa-distorção monoterminar de uma ampla classe de fontes Gaussianas não-estacionárias admite uma representação por preenchimento (*waterfilling*) reverso no domínio tempo-frequência [55]. Essa perspectiva estabelece uma ponte sólida entre a teoria de operadores e a teoria da informação clássica, substituindo operações matriciais de alta dimensão por integrais no plano tempo-frequência.

Apesar desses avanços, uma caracterização tempo-frequência comparável para problemas multiterminais com fontes Gaussianas não-estacionárias permaneceu, em grande parte, pouco desenvolvida. Em particular, pouco se sabe sobre como a taxa se comporta quando a fonte apresenta uma estrutura espectral suave (por exemplo, leis de potência tipo Sobolev), enquanto o ruído de observação varia suavemente no tempo. Nessas situações, nem a teoria clássica do problema do CEO vetorial nem os resultados para taxa-distorção monoterminar não-estacionária se aplicam diretamente.

5.2 Formulação do Problema

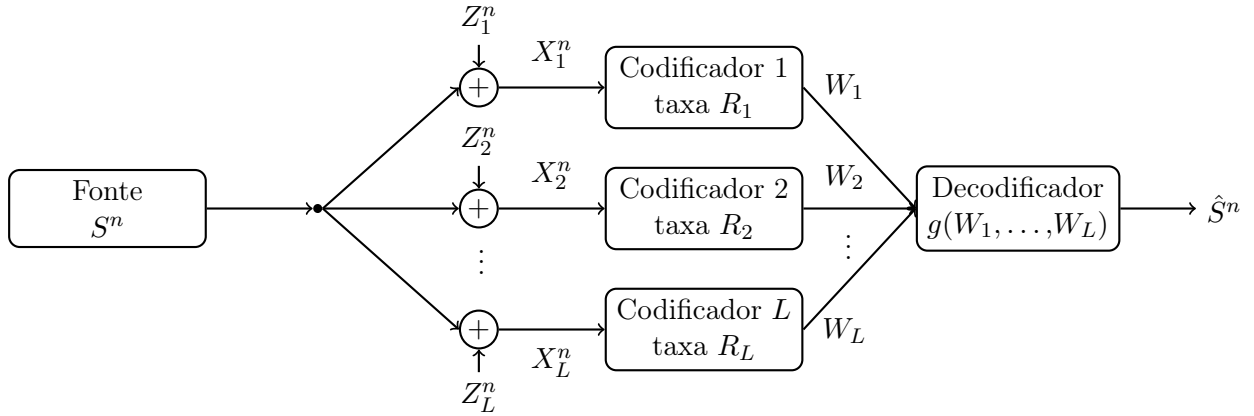


Figura 12: Codificação de fonte distribuída com ruídos aditivos independentes na observação de cada sensor.

O problema do CEO genérico é ilustrado na Figura 12 para o cenário no qual os ruídos dos sensores são aditivos. Vamos partir do princípio de que tanto a fonte quanto o ruído são Gaussianos e que a distorção é quadrática, mas, em geral, não-separável.

Mais precisamente, para um determinado tamanho de bloco n , seja a fonte $S^n \sim \mathcal{N}(0, \Sigma_S^{(n)})$ com covariância $\Sigma_S^{(n)} = T_n(s)$, uma matriz Toeplitz gerada por uma *Power Spectral Density* (PSD) real e limitada $s(\omega)$. Isso permite introduzir suposições realistas de suavidade na estrutura da fonte, úteis para imagens e campos escalares, por exemplo, ao modelar a PSD por um decaimento exponencial $s(\omega) = (1 + |\omega|^2)^{-\alpha}$.

Sejam as observações

$$X_\ell^n = S^n + Z_\ell^n, \quad \ell = 1, \dots, L, \quad (9)$$

onde $Z_\ell^n \sim \mathcal{N}(0, \Sigma_\ell^{(n)})$ são independentes ao longo dos terminais ℓ e suas covariâncias são diagonais $\Sigma_\ell^{(n)} = \text{diag}(\sigma_\ell^2(1/n), \dots, \sigma_\ell^2(1))$, permitindo levar em conta a variação da qualidade das observações ao longo do tempo. Supomos que os perfis de variância $\sigma_\ell^2(i/n)$ variam de forma suave no tempo, i.e., cada $\sigma_\ell^2 : [0, 1] \rightarrow (0, \infty)$ é Hölder-contínua com expoente γ_ℓ .

Finalmente, seja $A_n = T_n(a) \succ 0$ uma matriz Toeplitz gerada por um símbolo limitado $a(\omega)$, e defina-se a distorção quadrática ponderada $d(S^n, \hat{S}^n) = (S^n - \hat{S}^n)^T A_n (S^n - \hat{S}^n)$. No caso em que $a(\omega) = (1 + |\omega|^2)^\beta$, isso corresponde a um modelo de Sobolev de ordem β , o que permite uma ênfase dependente da frequência na precisão da reconstrução.

Defina-se um código de fonte como um conjunto de funções codificadoras

$$f_\ell : \mathbb{R}^n \rightarrow \{1, \dots, e^{nR_\ell}\}, \quad \ell = 1, \dots, L,$$

que mapeiam cada observação a uma mensagem $W_\ell = f_\ell(X_\ell^n)$ a taxa R_ℓ , e uma função decodificadora

$$g : \prod_{\ell=1}^L \{1, \dots, e^{nR_\ell}\} \rightarrow \mathbb{R}^n$$

que produz uma estimativa $\hat{S}^n = g(W_1, \dots, W_L)$.

A falta de uma autobase comum entre a covariância da fonte e as covariâncias dos ruídos significa que a diagonalização simultânea para desacoplar o problema em n problemas escalares não é viável. Ao invés disso, considere uma aproximação das variações temporais por funções degrau. Fixe um número B de blocos e particione $[0,1]$ em intervalos $I_b \equiv ((b-1)/B, b/B]$, cada um com $m \equiv n/B$ elementos. Defina as aproximações por degrau $\sigma_{\ell,B}(t) = \sigma_\ell(b/B)$ para $t \in I_b$, e seja

$$\Sigma_{\ell,B}^{(n)} = \text{diag}(\sigma_{\ell,B}^2(1/n), \dots, \sigma_{\ell,B}^2(1)). \quad (10)$$

Considere também a construção de um intervalo de guarda de h amostras ao final de cada bloco, $h < m$. Defina o conjunto de amostras mantidas $J \equiv \bigcup_{b=1}^B \{(b-1)m+1, \dots, bm-h\}$ e seja $\bigoplus_{b=1}^B T_{m-h}(\cdot)$ a soma direta de B cópias de $T_{m-h}(\cdot)$, ou seja, a matriz diagonal por blocos $(n-Bh) \times (n-Bh)$ formada pela repetição de $T_{m-h}(\cdot)$ B vezes. Precisaremos dos seguintes lemas de equivalência assintótica e de estabilidade de Lipschitz, a fim de justificar a solução proposta no domínio tempo-frequência.

Lema 1 (equivalência assintótica). *Para uma função f com coeficientes de Fourier absolutamente somáveis $\hat{f}(k)$ e para o conjunto J descrito acima,*

$$\|(T_n(f))_J - \bigoplus_{b=1}^B T_{m-h}(f)\| \leq \sum_{|k|>h} |\hat{f}(k)|, \quad (11)$$

onde o subscrito J indica uma submatriz principal.

Além disso, para σ_ℓ^2 Hölder-contínua com expoente γ_ℓ ,

$$\|\Sigma_\ell^{(n)} - \Sigma_{\ell,B}^{(n)}\| \leq c_\ell B^{-\gamma_\ell}. \quad (12)$$

Seja $\theta^n \equiv (S^n, X_1^n, \dots, X_L^n)$ o vetor Gaussiano de dimensão $(L+1)n$ com covariância $\Sigma^{(n)}$. Note que $\|\Sigma^{(n)}\| \leq (L+1)\|\Sigma_S^{(n)}\| + \max_\ell \|\Sigma_\ell^{(n)}\|$. O problema do CEO Gaussiano quadrático é, assim, completamente descrito pelas matrizes Σ e A , as quais definem a distribuição conjunta da fonte e a métrica de distorção, respectivamente. Nosso objetivo é derivar uma aproximação assintótica para a função de taxa-soma-distorção

$$R_n(D \mid \Sigma, A) \equiv \inf \left\{ \sum_{\ell=1}^L R_\ell : \exists (f_1, \dots, f_L, g) \text{ t.q. } \mathbb{E}[d_A(S^n, \hat{S}^n)] \leq nD \right\}. \quad (13)$$

Lema 2 (estabilidade Lipschitz). Se $|\mathbb{E}_P[d] - \mathbb{E}_Q[d]| \leq n\Delta_{P,Q}$, então

$$|R_n(D | P) - R_n(D | Q)| \leq c_R \Delta_{P,Q} \quad (14)$$

Dois corolários imediatos do Lema 2 são:

Corolário 1.

$$|R_n(D | \Sigma, A) - R_n(D | \bar{\Sigma}, A)| \leq c_{1,\Sigma} \|\Sigma - \bar{\Sigma}\| + c_{2,\Sigma} \|\Sigma - \bar{\Sigma}\|^2 \quad (15)$$

Corolário 2.

$$|R_n(D | \Sigma, A) - R_n(D | \Sigma, \bar{A})| \leq c_A \|A - \bar{A}\| \quad (16)$$

No modelo aproximado por truncamentos em matrizes diagonais por blocos da covariância da fonte e dos pesos da distorção, e com perfis de ruído constantes por bloco, o problema se decompõe em B blocos estacionários independentes, descritos pelas matrizes $\bar{\Sigma}_J$ e $\bar{A}_J$. Para o bloco b , com comprimento $m_b \equiv m - h$ e variâncias de ruído constantes $\sigma_\ell^2(b/B)$, a solução para cada bloco passa a ser a conhecida solução de Berger-Tung vetorial $R_{m_b}^{\text{BT}}$, e a decomposição é

$$R_n^{(J)}(D | \bar{\Sigma}_J, \bar{A}_J) = \inf_{\{D_b\}: \frac{1}{B} \sum_b D_b = D} \frac{m_b}{n} \sum_b R_{m_b}^{\text{BT}}(D_b). \quad (17)$$

Pelo teorema de Szegő, a taxa-soma estacionária de Berger-Tung para cada bloco individual, conforme $m_b \rightarrow \infty$, torna-se

$$R_{\text{st}}(D_b) = \inf_{d_b(\cdot), r_{\ell,b}(\cdot)} \frac{1}{2\pi} \int_{-\pi}^{\pi} \left[\frac{1}{2} \log \left(\frac{s(\omega)}{d_b(\omega)} \right) + \sum_{\ell} r_{\ell,b}(\omega) \right] d\omega, \quad (18)$$

onde $r_{\ell,b}(\omega)$ são as taxas do canal de teste para cada codificador, sujeitas às usuais restrições de Berger-Tung

$$\sum_{\ell} \frac{1 - e^{-2r_{\ell,b}(\omega)}}{\sigma_\ell^2(b/B)} = \frac{1}{d_b(\omega)} - \frac{1}{s(\omega)} \quad (19)$$

$$\frac{1}{2\pi} \int a(\omega) d_b(\omega) d\omega \leq D_b. \quad (20)$$

5.3 Resultados

Nesta seção apresentaremos nossos principais resultados. Antes de apresentar o primeiro resultado, observamos que a soma de Riemann em (17) se transforma em uma integral tempo-frequência à medida que $B \rightarrow \infty$. Defina $d(t, \omega) = d_b(\omega)$ para $t \in I_b$. Então

$$\lim_{B \rightarrow \infty} \inf_{\{D_b\}} \frac{m_b}{n} \sum_b R_{\text{st}}(D_b) = R_{\infty}(D), \quad (21)$$

onde definimos o limite variacional em tempo-frequência $R_{\infty}(D)$ por

$$R_{\infty}(D) \equiv \inf_{d(\cdot, \cdot), r_{\ell}(\cdot, \cdot)} \frac{1}{2\pi} \int \int \left[\frac{1}{2} \log \left(\frac{s(\omega)}{d(t, \omega)} \right) + \sum_{\ell} r_{\ell}(t, \omega) \right] dt d\omega \quad (22)$$

$$\text{s.a.} \quad \sum_{\ell} \frac{1 - e^{-2r_{\ell}(t, \omega)}}{\sigma_\ell^2(t)} = \frac{1}{d(t, \omega)} - \frac{1}{s(\omega)} \quad (23)$$

$$\frac{1}{2\pi} \int \int a(\omega) d(t, \omega) dt d\omega \leq D. \quad (24)$$

Teorema 1 (caracterização de primeira ordem). *Sob as premissas do modelo apresentado, para todo D no interior do intervalo de distorção factível,*

$$\lim_{n \rightarrow \infty} R_n(D) = R_\infty(D). \quad (25)$$

Teorema 2 (condições de *Karush–Kuhn–Tucker* (KKT)). *O problema de otimização com restrições em (22)-(24) admite uma solução por waterfilling. Existem $\mu > 0$ e $\nu(t, \omega)$ tais que o ótimo (d^*, r_ℓ^*) satisfaz*

$$r_\ell^*(t, \omega) = \frac{1}{2} \log^+ \left(\frac{\nu(t, \omega)}{\sigma_\ell^2(t)} \right) \quad (26)$$

$$\frac{1}{d^*(t, \omega)} - \frac{1}{s(\omega)} = \sum_{\ell=1}^L \left(\frac{1}{\sigma_\ell^2(t)} - \frac{1}{\nu(t, \omega)} \right)^+ \quad (27)$$

$$\nu(t, \omega) = 2\mu a(\omega) d^{*2}(t, \omega) - d^*(t, \omega) \quad (28)$$

$$D = \frac{1}{2\pi} \int_{-\pi}^{\pi} \int_0^1 a(\omega) d(t, \omega) dt d\omega. \quad (29)$$

O conjunto factível pontual é definido por um intervalo de valores de distorção:

$$\left(\frac{1}{s(\omega)} + \sum_{\ell} \frac{1}{\sigma_\ell^2(t)} \right)^{-1} < d(t, \omega) \leq s(\omega) \quad (30)$$

Usando (27)-(28) para induzir uma solução pontual para $d^*(t, \omega)$ parametrizada por μ , é possível aplicar a restrição global (29) para encontrar uma solução única para μ .

5.4 Conclusão

Desenvolvemos uma caracterização de primeira ordem no domínio tempo-frequência do problema do CEO Gaussiano não-estacionário sob distorção quadrática genérica. Sob estrutura de fonte Toeplitz/Sobolev e ruídos de observação heterocedásticos com variação suave, o problema do CEO vetorial inerentemente não-comutativo admite uma redução assintoticamente tratável a um conjunto contínuo de subproblemas do CEO escalares locais acoplados por meio de uma restrição de distorção global. A expressão integral resultante pode ser interpretada como uma regra de preenchimento (*waterfilling*) no domínio tempo-frequência sobre um campo de canais locais, com um único multiplicador de Lagrange governando a alocação da distorção ao longo do tempo e da frequência.

Além de sua contribuição teórica, o modelo sugere uma estratégia de codificação conceitual baseada na codificação local de Berger-Tung em blocos de tempo-frequência, binagem Slepian-Wolf, e *waterfilling* global. Trabalhos futuros poderão estender essa abordagem a modelos de observação lineares ou a fontes com memória temporal, além de quantificar os efeitos de tamanhos finitos de blocos e as caracterizações assintóticas de segunda ordem em contextos de problemas localmente estacionários.

6 Rastreamento Remoto de Processos Baseado em Comunicação Orientada a Tarefas e Predição

Isaac Andrés Balbuca Guachizaca e Paulo Cardieri
UNICAMP

6.1 Introdução

Ao longo do tempo, as telecomunicações avançaram em paralelo com a forma como as pessoas interagem com a informação. Inicialmente, as redes foram projetadas para a transmissão de voz. Posteriormente, evoluíram para transportar conteúdo multimídia e conectar uma ampla variedade de dispositivos, criando uma sociedade hiperconectada. O desenvolvimento das comunicações sem fio, impulsionado pelas redes 6G e pela proliferação de sistemas IoT e dispositivos autônomos, exige mudanças na forma como as redes operam, pois deixam de ser meros canais de transporte de dados para se tornarem plataformas de inteligência interconectada [56].

O paradigma de comunicação tradicional enfrenta limitações diante de aplicações emergentes que impõem requisitos críticos de latência e confiabilidade, além de severas restrições de energia. Tais demandas levam os sistemas convencionais a operar próximos aos limites de capacidade física definidos por Shannon [57].

Nesse contexto, torna-se necessário revisitar os fundamentos estabelecidos por Shannon e Weaver, cuja Teoria Matemática da Comunicação definiu três problemas fundamentais relacionados ao processo de comunicação: (i) o *problema técnico*, que trata da precisão na transmissão de símbolos; (ii) o *problema semântico*, focado em transmitir o significado da mensagem; e (iii) o *problema da efetividade*, que avalia o impacto da mensagem na realização de uma tarefa [58].

Historicamente, os sistemas de comunicação tradicionais concentraram-se no problema técnico. Contudo, os avanços em inteligência artificial e em poder computacional deram origem a duas novas áreas de pesquisa: a **Comunicação Semântica**, que explora o problema semântico, e as **Comunicações Orientadas a Tarefas**, que se concentram no problema da efetividade [59].

Na comunicação orientada a tarefas, o objetivo é otimizar as estratégias de transmissão para garantir a conclusão bem-sucedida de uma tarefa predefinida com o menor uso possível de recursos de comunicação. Ao contrário dos sistemas convencionais, que priorizam a entrega de toda a informação bruta, independentemente de sua relevância, os sistemas orientados a tarefas transmitem apenas os dados mais relevantes para a tarefa que requereu a comunicação. Essa abordagem seletiva, porém, requer cooperação entre os terminais de origem e de destino, bem como conhecimento dos objetivos da tarefa pelo transmissor.

Políticas de transmissão orientadas a tarefas foram introduzidas para explorar a relevância semântica das mensagens em diversos problemas. Neste trabalho, focamos no problema de rastreamento remoto de processos, no qual um sinal gerado por um terminal fonte deve ser reconstruído no terminal de destino. Pappas et al. [60] e, posteriormente, Salimnejad et al. [61] investigaram esse problema considerando um processo de Markov. Mais recentemente, estudos sobre sistemas IoT com coleta de energia obtiveram políticas ótimas baseadas em limiares que consideram conjuntamente o estado da fonte, a idade da informação e os níveis da bateria [62].

Este estudo investiga uma estratégia de transmissão para o problema de rastreamento remoto sob restrições de energia. Consideramos o mesmo cenário adotado por Pappas et al. em [60], ou seja, o processo a ser rastreado é modelado como um processo de Markov X_t e, por meio de mensagens de controle ACK e NACK, o transmissor tem conhecimento do processo recons-

truído no destino $\hat{X}_t$. Combinamos a política *Semantics-Aware*, proposta por Pappas et al. em [60], com uma técnica preditiva para reduzir o número de transmissões e, consequentemente, o consumo de energia do transmissor. Demonstramos que, sob restrições severas de energia, a abordagem proposta é capaz de preservar energia para transmissões futuras, reduzindo, assim, o erro de reconstrução do sinal no destino.

O restante deste capítulo está organizado da seguinte forma. A Seção 6.2 descreve o modelo de sistema adotado na análise das políticas de transmissão. A Seção 6.3 apresenta as políticas de transmissão avaliadas neste trabalho: a política de referência *Semantics-Aware*, a política Preditiva proposta e a política Preditiva Ideal. Os resultados numéricos são apresentados e discutidos na Seção 6.4, e as conclusões são apresentadas na Seção 6.5.

6.2 Modelo de Sistema

Consideramos o sistema ilustrado na Fig. 13, no qual um transmissor observa um processo estocástico X_t nos instantes de tempo $t \in \{0, 1, \dots\}$ e envia atualizações de estado a um receptor remoto seguindo a política de transmissão selecionada. O receptor reconstrói o processo fonte,

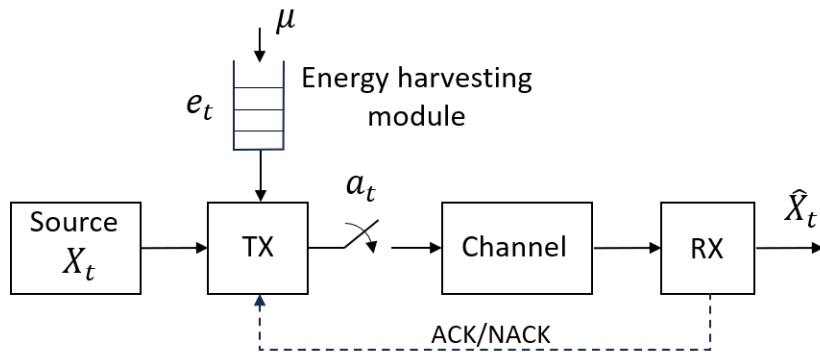


Figura 13: Modelo de sistema para comunicação semântica orientada a tarefas com coleta de energia.

denotado por $\hat{X}_t$.

O processo X_t é modelado como um processo de Markov de nascimento e morte (BDMP) com N estados, definido sobre um espaço de estados finito $\mathcal{S} = \{0, 1, \dots, N-1\}$. As transições de estado são restritas ao estado atual ou aos seus vizinhos imediatos. Especificamente, a fonte transita do estado i para $i+1$ com probabilidade de nascimento p e para $i-1$ com probabilidade de morte q , permanecendo no estado atual com probabilidade $1-p-q$, sujeita às condições de contorno.

O funcionamento do sistema segue o modelo adotado por Pappas et al. em [60], conforme descrito a seguir. A cada *time slot* t , o transmissor decide se deve enviar uma atualização de acordo com a variável de decisão binária $a_t \in \{0, 1\}$, em que $a_t = 1$ denota uma tentativa de transmissão e, caso contrário, $a_t = 0$. O valor de a_t é determinado pela estratégia de transmissão adotada. As transmissões ocorrem por um canal com probabilidade de sucesso p_s , e o receptor envia um reconhecimento positivo (ACK) para uma transmissão bem-sucedida ou um reconhecimento negativo (NACK) para uma falha, assumidos como entregues instantaneamente e sem erros. Por meio desse retorno, o transmissor mantém conhecimento perfeito do estado reconstruído $\hat{X}_t$ no receptor em qualquer instante. Consequentemente, o transmissor pode comparar o estado atual da fonte X_t com o estado reconstruído no receptor $\hat{X}_t$ em qualquer

time slot. Isso permite detectar discrepâncias e tomar decisões com base na sua política de transmissão.

O transmissor é alimentado por um módulo de Coleta de Energia (*energy harvesting*) equipado com uma bateria de capacidade finita, denotada por $E_{\max}$. De forma análoga a [63], assumimos que as chegadas de energia seguem um processo de Bernoulli com parâmetro μ . Definimos um indicador binário $u_t \in \{0,1\}$ para representar a chegada de energia no início do *time slot* t , com $\Pr[u_t = 1] = \mu$. Quando $u_t = 1$, uma unidade de energia é coletada e adicionada à bateria. Consequentemente, o nível da bateria e_t evolui de acordo com a seguinte dinâmica:

$$e_{t+1} = \min\{e_t + u_t - a_t c_t, E_{\max}\}, \quad (31)$$

em que c_t denota o custo de energia associado a uma tentativa de transmissão.

Uma transmissão é permitida no *time slot* t somente se a bateria possuir energia suficiente, o que impõe a restrição de energia $e_t - a_t c_t \geq 0$, $\forall t$.

A energia consumida pelo transmissor depende do número de tentativas de transmissão, o qual é determinado tanto pela estratégia de transmissão adotada quanto pelas condições do canal, uma vez que transmissões malsucedidas podem desencadear comportamentos diferentes no transmissor. Portanto, para possibilitar uma comparação justa entre as diferentes estratégias de transmissão, o custo de energia deve ser adequadamente associado à probabilidade de sucesso de transmissão desejada. Para fins de tratabilidade, consideramos um canal sem fio AWGN com modulação BPSK. Sob essa hipótese, a energia necessária para a transmissão bem-sucedida de um pacote contendo L bits com probabilidade p_s é dada por⁴

$$E = \frac{L N_0}{2} [Q^{-1}(1 - p_s^{1/L})]^2, \quad (32)$$

em que N_0 denota a densidade espectral de potência do ruído e $Q(\cdot)$ denota a função Q gaussiana.

6.3 Políticas de Transmissão

Nesta seção, descrevemos as políticas de transmissão avaliadas neste estudo. Primeiramente, apresentamos a política *Semantics-Aware*, que serve como referência de desempenho estabelecida na literatura. Em seguida, introduzimos a política Preditiva proposta e sua versão ideal.

Política *Semantics-Aware* (SAP)

A política *Semantics-Aware*, proposta por Pappas e Kountouris [60], baseia-se em duas hipóteses: (a) o transmissor possui conhecimento perfeito do estado reconstruído $\hat{X}_t$, e (b) o receptor atualiza seu estado reconstruído sempre que uma mensagem contendo uma nova atualização de estado é recebida com sucesso. Caso contrário, o receptor mantém seu estado reconstruído anterior.

Considere o funcionamento do sistema no *time slot* $t + 1$. Se o sistema está sincronizado, ou seja, $X_{t+1} = \hat{X}_t$, nenhuma transmissão é necessária, pois, na ausência de uma nova atualização, o receptor simplesmente mantém sua estimativa atual, isto é, $\hat{X}_{t+1} \leftarrow \hat{X}_t$. Por outro lado, se existe uma discrepância, ou seja, $X_{t+1} \neq \hat{X}_t$, o transmissor envia um pacote de atualização ao receptor. Se a transmissão for bem-sucedida, o receptor atualiza sua estimativa para $\hat{X}_{t+1} =$

⁴Assumimos que um pacote é recebido com sucesso somente se todos os L bits forem corretamente decodificados.

X_{t+1} . Caso contrário, se a transmissão falhar, nenhuma retransmissão é tentada naquele *time slot*, e o receptor mantém sua estimativa anterior, ou seja, $\hat{X}_{t+1} = \hat{X}_t$. Note que a política *Semantics-Aware* é orientada a tarefas, pois transmite apenas quando o estado do receptor diverge do estado real da fonte, garantindo que cada transmissão contribua diretamente para a redução do erro de reconstrução.

Política Preditiva (PP)

Na política *Semantics-Aware*, uma transmissão é acionada sempre que o sistema fica fora de sincronismo. Embora essa estratégia seja eficaz para manter o receptor atualizado, pode acarretar um elevado consumo de energia. Motivados por essa limitação, e com o objetivo de reduzir simultaneamente o consumo de energia e manter o erro de reconstrução em níveis aceitáveis, propomos a *Política Preditiva*, que pode ser interpretada como uma extensão da política *Semantics-Aware*.

A política Preditiva proposta opera da seguinte forma. Quando o sistema está sincronizado ($X_t = \hat{X}_t$), nenhuma atualização é necessária. No entanto, se existe uma discrepância, um preditor estima o próximo estado da fonte, denotado por $\tilde{X}_{t+1}$, e o compara com o estado atual no receptor $\hat{X}_t$. Se a predição indica que o sistema se ressincronizará naturalmente (ou seja, $\tilde{X}_{t+1} = \hat{X}_t$), nenhuma transmissão é realizada, economizando energia ao custo de um aumento temporário no erro de reconstrução, já que o erro no instante t não é corrigido. Por outro lado, se a discrepância tende a persistir, ou seja, $\tilde{X}_{t+1} \neq \hat{X}_t$, uma transmissão é acionada para restaurar a sincronização no receptor.

Qualquer técnica de predição pode ser empregada para estimar $\tilde{X}_{t+1}$. Neste trabalho, adotamos uma abordagem probabilística simples na qual o transmissor estima recursivamente as probabilidades de transição da fonte, denotadas por $\hat{p}$ e $\hat{q}$, a partir dos eventos de nascimento e morte observados. Com base nessas probabilidades de transição estimadas, o próximo estado predito $\tilde{X}_{t+1}$ é escolhido como o estado futuro mais provável da fonte.

Política Preditiva Ideal (IPP)

Para estabelecer um limite superior de desempenho para a política Preditiva, introduzimos a *Política Preditiva Ideal*. Nessa política, o transmissor possui conhecimento perfeito do próximo estado da fonte X_{t+1} no momento da decisão, eliminando qualquer erro de predição. Embora fisicamente irrealizável, essa política serve como referência que isola a contribuição dos erros de predição no desempenho global de reconstrução. O erro residual da IPP decorre apenas de perdas no canal e de transmissões intencionalmente suprimidas.

6.4 Resultados Numéricos

O desempenho das políticas de transmissão apresentadas na seção anterior é avaliado por meio de simulações em diferentes cenários operacionais representativos.

6.4.1 Configuração das Simulações e Cenários

A fonte de informação é modelada como um BDMP com $N = 10$ estados, considerando dois cenários distintos em relação à variabilidade da fonte:

- **Evento Raro:** Representa uma fonte de variação lenta, caracterizada pelas probabilidades de transição $p = 0.01$ e $q = 0.50$.

- **Fonte Rápida:** Representa uma fonte de variação rápida, caracterizada por $p = 0.20$ e $q = 0.70$.

Ambos os cenários são avaliados sob diferentes restrições de energia, variando a capacidade da bateria ($E_{\max} \in \{2, 10\}$) e a probabilidade μ de que uma unidade de energia seja coletada durante um *time slot*. Esses cenários com restrições de energia permitem avaliar os potenciais benefícios das políticas Preditivas, que buscam preservar recursos de energia ao custo de um aumento temporário no erro de reconstrução.

Duas condições de canal representativas são consideradas: $p_s = 0.5$, correspondente a um *canal não confiável*, e $p_s = 0.9$, correspondente a um *canal confiável*.

Quanto ao custo de transmissão c_t , assumimos que atingir uma probabilidade de sucesso de transmissão de $p_s = 0.5$ requer $c_t = 1$ unidade de energia. Utilizando a modelagem apresentada na Seção 6.2 (Eq. (32)), ao aumentar-se a confiabilidade do canal para $p_s = 0.9$, o transmissor requer aproximadamente $c_t = 1.56$ unidade de energia por transmissão, assumindo um tamanho de pacote de $L = 100$ bits.

Em todas as figuras apresentadas a seguir, os seguintes acrônimos são utilizados: SAP = Política *Semantics-Aware*, IPP = Política Preditiva Ideal, e PP = Política Preditiva.

Nas Seções 6.4.3 e 6.4.4, comparamos o desempenho das políticas *Semantics-Aware* e Preditiva. Os resultados obtidos com a Política Preditiva Ideal são discutidos na Seção 6.4.5, juntamente com discussões gerais.

6.4.2 Métricas de Desempenho

Para quantificar a eficiência das políticas, consideramos as seguintes métricas de desempenho.

Erro de reconstrução

Esta métrica, conforme definida em [60], mede a discrepância entre o processo fonte original X_t e o estado reconstruído no receptor $\hat{X}_t$. O erro de reconstrução médio temporal ao longo de um intervalo de observação grande T (expresso em número de *time slots*) é dado por

$$\bar{E} = \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=1}^T \mathbb{1}(X_t \neq \hat{X}_t), \quad (33)$$

em que $\mathbb{1}(\cdot)$ é a função indicadora.

Taxa de consumo de energia

Para avaliar a eficiência energética das políticas de transmissão, utilizamos a *taxa de consumo de energia*, definida como

$$\bar{C} = \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=1}^T a_t c_t, \quad (34)$$

em que $a_t \in \{0, 1\}$ é a decisão binária de transmissão no *time slot* t e c_t é o custo de energia associado a uma tentativa de transmissão, conforme introduzido na Seção 6.2. Essa métrica representa a energia média consumida por *time slot* ao longo de um grande intervalo de observação T .

Ambas as métricas foram avaliadas tomando-se a média sobre $T = 10^5$ *time slots*.

6.4.3 Resultados para o Cenário de Evento Raro

Canal não confiável ($p_s = 0.5$)

A Fig. 14a apresenta o erro de reconstrução alcançado pelas três políticas em função da probabilidade de coleta de energia, para diferentes valores de capacidade da bateria.

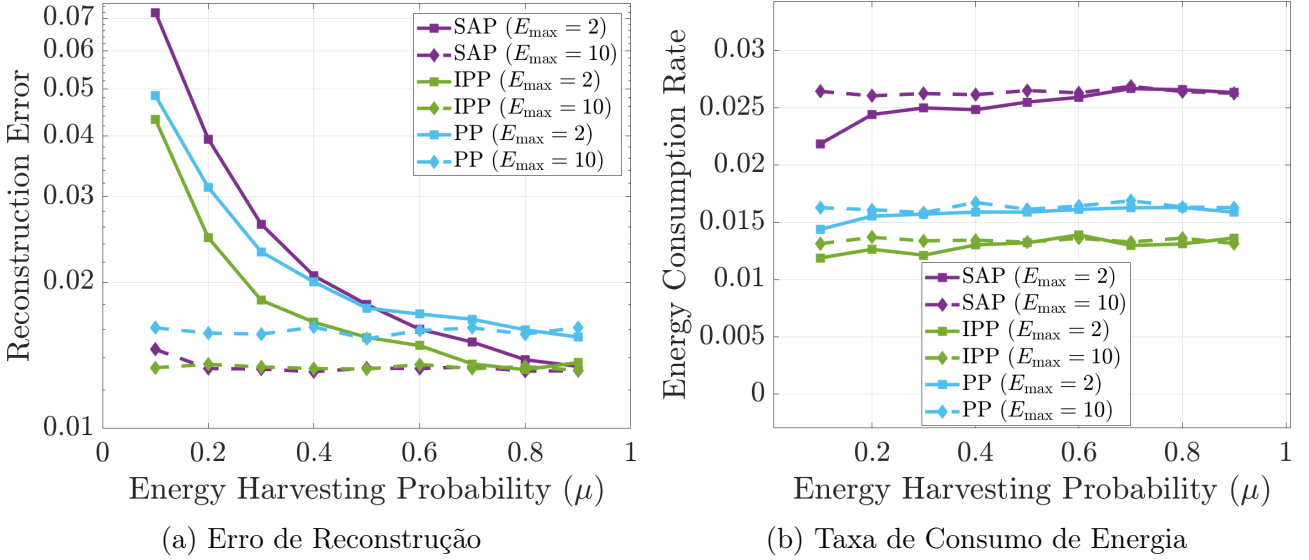


Figura 14: Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$, no cenário de Evento Raro com $p_s = 0.50$.

Sob restrições severas de energia ($E_{\max} = 2$ e $\mu < 0.5$), a política *Semantics-Aware* apresenta o maior erro de reconstrução. Isso ocorre porque, para cada discrepância detectada, o transmissor tenta uma correção por um canal que falha metade das vezes, esgotando rapidamente a bateria e impossibilitando transmissões futuras. Por outro lado, a Política Preditiva, ao antecipar as autocorreções, tolera erros transitórios e aciona transmissões apenas quando a discrepância tende a persistir. Dessa forma, a energia é preservada para transmissões futuras, levando a um desempenho superior ao da política *Semantics-Aware*. No entanto, para $\mu > 0.5$, a bateria é recarregada com frequência suficiente para que a política *Semantics-Aware* corrija as discrepâncias de forma eficaz, superando assim a Política Preditiva.

Para $E_{\max} = 10$, a maior capacidade da bateria permite que a política *Semantics-Aware* transmita sempre que necessário, atingindo assim o menor erro de reconstrução.

A Fig. 14b mostra que, como esperado, as políticas Preditivas são mais eficientes em termos de energia. Ao estimar o estado futuro da fonte, essas políticas acionam transmissões apenas quando a discrepância tende a persistir. Como resultado, o consumo de energia é significativamente menor.

De modo geral, embora a política *Semantics-Aware* possa alcançar menor erro de reconstrução quando a energia é abundante, ela opera com um maior consumo de energia, o que pode impedir transmissões e elevar o erro de reconstrução se a capacidade da bateria for pequena. Por outro lado, a Política Preditiva mantém maior eficiência energética aceitando um pequeno erro residual.

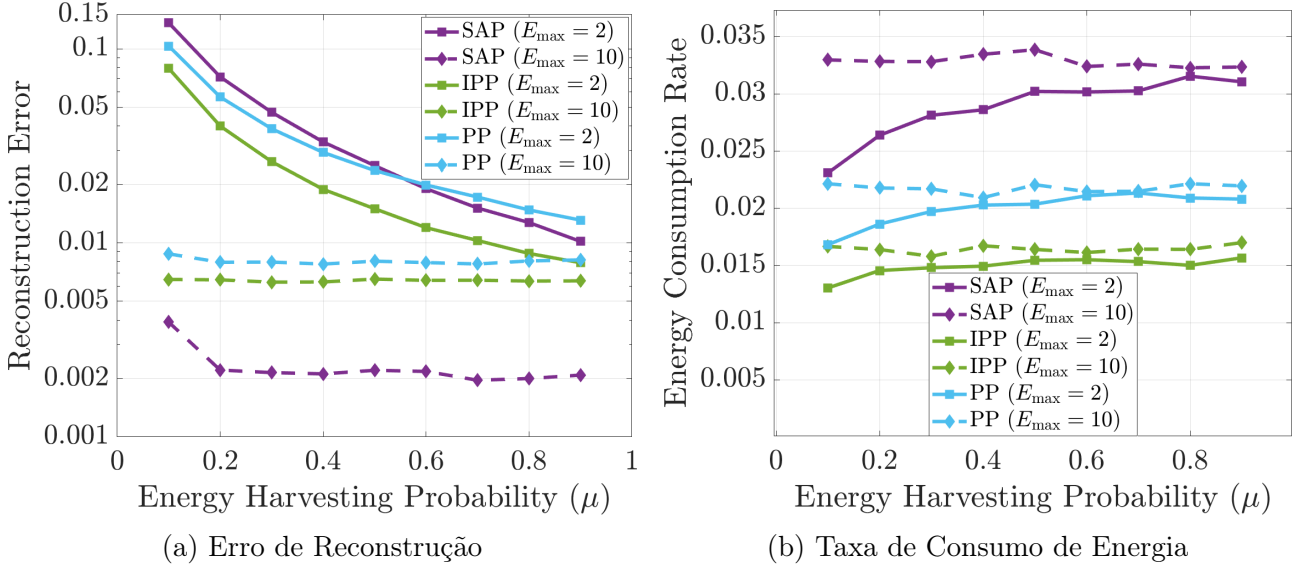


Figura 15: Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$, no cenário de *Evento Raro* com $p_s = 0.90$.

Canal confiável ($p_s = 0.9$)

Operar em um canal confiável requer um custo de energia por transmissão maior ($c_t \approx 1.56$). Consequentemente, sob restrições severas de energia ($E_{\max} = 2$ e $\mu < 0.5$), uma única transmissão consome quase toda a capacidade da bateria, levando à degradação do desempenho de todas as políticas (ver Fig. 15a), especialmente da política *Semantics-Aware*. Por outro lado, a maior confiabilidade do canal aumenta a probabilidade de entrega bem-sucedida dos pacotes, mitigando parcialmente os efeitos adversos da capacidade limitada da bateria.

Para uma capacidade maior de bateria ($E_{\max} = 10$), a política *Semantics-Aware* dispõe de energia suficiente para transmitir quase todos os pacotes de atualização, como no cenário de canal não confiável. No caso de canal confiável, porém, ela alcança um erro de reconstrução substancialmente menor do que as políticas Preditivas, uma vez que 90% dos pacotes de atualização transmitidos são recebidos com sucesso em média.

Como no caso do canal não confiável, as políticas preditivas são mais eficientes em termos de energia, conforme mostrado na Fig. 15b.

6.4.4 Resultados para o Cenário de Fonte Rápida

Canal não confiável ($p_s = 0.5$)

Nesse cenário com uma fonte de variação rápida, a demanda por atualizações de estado torna-se significativamente maior, acentuando ainda mais os efeitos adversos das restrições de energia, conforme ilustrado na Fig. 16a, com erros de reconstrução variando de 0.2 a 0.5. Como anteriormente, sob restrições severas de energia ($E_{\max} = 2$ e $\mu < 0.4$), a política *Semantics-Aware* exibe o pior desempenho. No entanto, à medida que essas restrições são relaxadas e mais energia se torna disponível, a política *Semantics-Aware* alcança o menor erro de reconstrução.

Embora os erros de reconstrução resultantes sejam elevados demais para aplicações práticas, esses resultados destacam como a combinação de restrições severas de energia, baixa qualidade

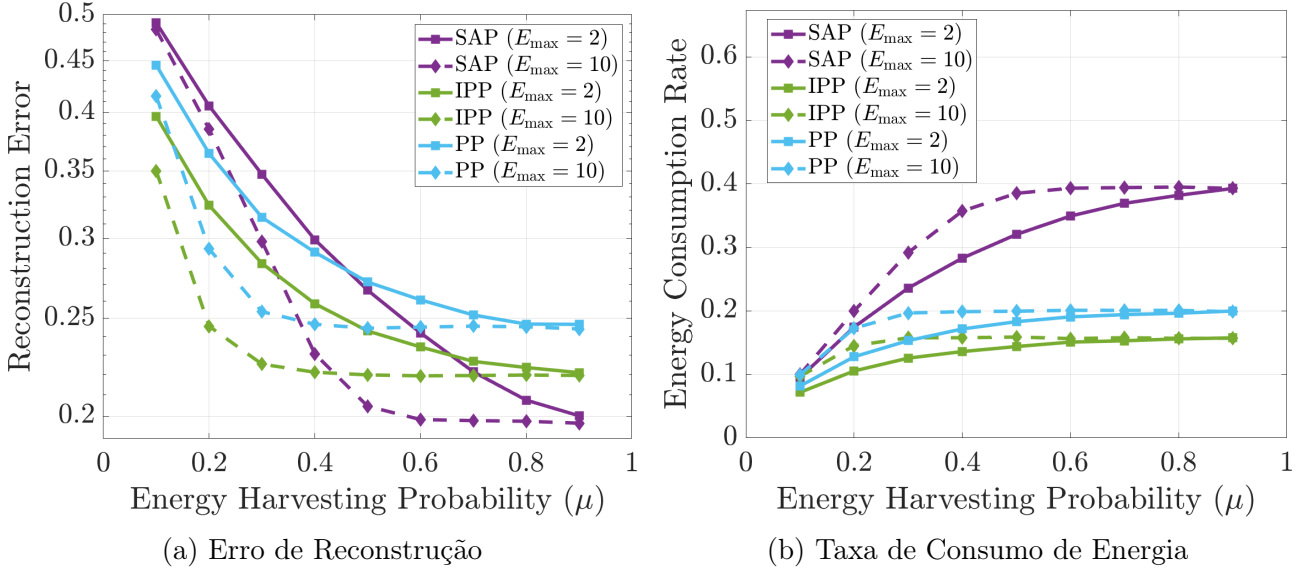


Figura 16: Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$ no cenário de Fonte Rápida com $p_s = 0.50$.

do canal e uma fonte de variação rápida exige políticas de transmissão adequadas para lidar com condições operacionais tão adversas.

Como esperado, todas as políticas consomem mais energia no cenário de Fonte Rápida do que no cenário de Evento Raro, como pode ser observado comparando as Figs. 14b e 16b. Além disso, no cenário de Fonte Rápida, a taxa de consumo de energia aumenta com a probabilidade de coleta de energia μ . Esse comportamento indica que a natureza altamente volátil da fonte esgota frequentemente a bateria, tornando a disponibilidade de energia fortemente dependente do processo de coleta.

Canal confiável ($p_s = 0.9$)

No caso do canal confiável, que requer um custo de transmissão maior ($c_t \approx 1.56$), as políticas apresentam comportamento semelhante ao observado no cenário de canal não confiável. No entanto, quando a energia é abundante, os resultados mostram que as políticas Preditivas atingem um patamar de erro significativamente maior do que o da política *Semantics-Aware*, mesmo no caso de predição ideal (ver Fig. 17). Isso indica que os erros intencionais introduzidos pelas políticas Preditivas representam uma parcela substancial do erro de reconstrução total quando a fonte varia rapidamente.

6.4.5 Discussão

Os resultados destacam os benefícios da política Preditiva em cenários com disponibilidade limitada de energia (baixo μ ou pequeno $E_{\max}$). Sob tais condições, a lógica de transmissão adotada pela política *Semantics-Aware*, que tenta corrigir imediatamente cada discrepância sem distinguir erros transitórios daqueles que naturalmente se autocorrigiriam no próximo *time slot*, pode levar ao esgotamento rápido da bateria. Como consequência, transmissões futuras tornam-se impossíveis, o que aumenta o erro médio de reconstrução. Em contraste, a política

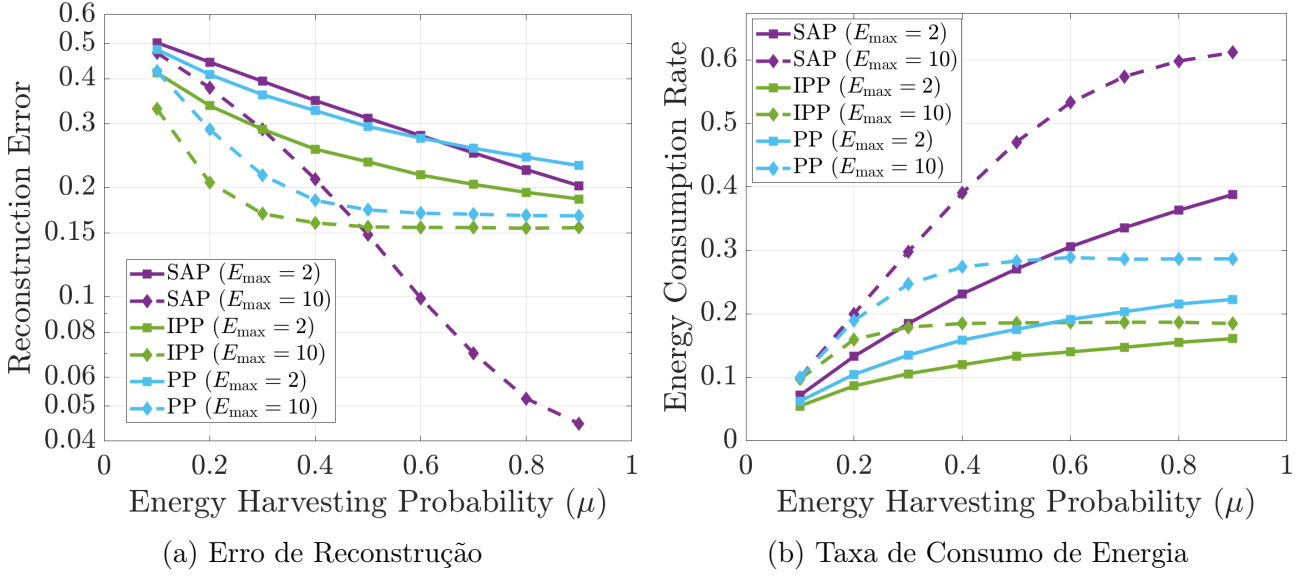


Figura 17: Erro de reconstrução (a) e taxa de consumo de energia (b) em função da probabilidade de coleta de energia, para capacidade da bateria $E_{\max} \in \{2, 10\}$ no cenário de Fonte Rápida com $p_s = 0.90$.

Preditiva busca preservar energia, mesmo ao custo de tolerar temporariamente um erro de reconstrução maior.

Os benefícios da política Preditiva também são observados em cenários com restrições de energia menos severas (alto μ ou grande $E_{\max}$), embora agora sob a perspectiva do consumo de energia. Nessas situações, a política Preditiva é capaz de alcançar níveis baixos de erro de reconstrução (abaixo de 1%, comparado a aproximadamente 0,2% alcançado pela política *Semantics-Aware*, conforme ilustrado na Fig. 15), com menor consumo de energia.

A diferença entre as políticas Preditiva e Preditiva Ideal fornece uma percepção adicional sobre os limites da abordagem preditiva. Uma vez que a política Preditiva Ideal elimina completamente os erros de predição, a diferença entre os erros de reconstrução alcançados pela Política Preditiva e pela Política Preditiva Ideal representa a degradação de desempenho causada exclusivamente pela incerteza estocástica do processo fonte, a qual não pode ser superada por nenhum preditor, independentemente da precisão com que p e q são estimados. A diferença é mais pronunciada no cenário de Fonte Rápida (ver Fig. 17a), em que as maiores probabilidades de transição aumentam a imprevisibilidade da fonte, tornando as predições menos confiáveis. Por outro lado, no cenário de Evento Raro, a fonte evolui lentamente e a diferença se reduz consideravelmente (ver Fig. 15a), indicando que a política Preditiva opera mais próxima do limite ideal quando a dinâmica da fonte é mais previsível. Essas observações sugerem que a política Preditiva é mais útil em um regime onde a fonte é suficientemente lenta para que as predições sejam confiáveis.

6.5 Conclusões e Trabalhos Futuros

Este estudo investigou o problema de rastreamento remoto de uma fonte markoviana sob restrições de energia, propondo uma Política Preditiva que equilibra o erro de reconstrução e o consumo de energia. Ao prever o estado da fonte, a política antecipa a autocorreção natural da fonte, preservando energia para erros persistentes em vez de esgotá-la em discrepâncias

transitórias, ao contrário da política *Semantics-Aware*. Os resultados revelam que, embora a política *Semantics-Aware* minimize o erro sob alta disponibilidade de energia, a política Preditiva oferece uma estratégia mais sustentável, alcançando economias substanciais de energia ao aceitar um erro residual.

Trabalhos futuros incluem a formulação do problema de escalonamento de transmissões como um Processo de Decisão de Markov para derivar uma política ótima que minimize conjuntamente o erro de reconstrução e o consumo de energia, incorporando o princípio de supressão de transmissões quando a ressincronização espontânea é antecipada.

7 Novo Protocolo de Acesso Aleatório baseado em Age of Information

Plínio Santini Dester, Richard Demo Souza, Paulo Cardieri
UNICAMP

7.1 Introdução

Com a crescente dependência de tomadas de decisão baseadas em dados em sistemas de larga escala e aplicações de IoT, a gestão do acesso ao meio torna-se um desafio central na arquitetura de redes 6G [64]. Para aplicações nas quais a entrega pontual de atualizações é crítica, a métrica AoI tem se consolidado como a mais amplamente adotada para quantificar o frescor da informação na rede [65], [66]. Diante de cenários de conectividade massiva de MTC, protocolos de acesso aleatório descentralizados oferecem soluções promissoras, pois viabilizam o acesso ao canal sem os elevados custos de sinalização e negociação de escalonamento [64].⁵

Uma abordagem eficiente e consolidada na literatura para incorporar a consciência de idade no acesso ao meio é o *Age-Threshold Slotted ALOHA* (TSA), concebido para permitir que um nó acesse o canal apenas quando a idade de sua atualização supera um limiar pré-definido [69]. Como contribuição para a concepção da rede de acesso neste projeto, propõe-se a modificação denominada 1-persistent TSA, que constitui um refinamento dessa dinâmica, na qual o nó transmite com probabilidade 1 exatamente no instante em que atinge seu limiar de idade.

Sob o modelo de tráfego *generate-at-will* (que foca na geração sob demanda), rede homogênea e canal sem apagamento, esse esquema atinge um desempenho estacionário teoricamente equivalente ao de um sistema *Time Division Multiple Access* (TDMA) perfeitamente coordenado. Ajustando-se o limiar de idade, obtém-se um compromisso entre a duração do regime transiente e o desempenho no estado estacionário. Nesse regime, o 1-persistent TSA supera o limite do TSA original (que fica em torno de 36,8% de throughput e 1,42 de AoI média normalizada [69]), atingindo um throughput de 50% e uma média de AoI normalizada igual a 1.

Em arquiteturas 6G, no entanto, é comum que os nós possuam requisitos distintos de frescor de informação, motivando o uso de restrições e limiares de AoI desiguais em redes heterogêneas [66]. A análise do 1-persistent TSA neste cenário revela uma propriedade fundamental do sistema: o protocolo converge *almost surely* para um regime estacionário determinístico e livre de colisões se, e somente se, os limiares de idade formarem uma *sequência harmônica*. Nesse regime, o throughput e as métricas de AoI podem ser expressos em forma fechada e tornam-se independentes da probabilidade de retransmissão adotada fora do limiar. A escolha de uma sequência harmônica permite que os períodos de transmissão se alinhem de tal modo que formem padrões periódicos não interferentes.

Para a concepção da rede de acesso 6G voltada a dispositivos IoT e MTC, o uso de limiares de idade para o controle de acesso, em especial sob o comportamento 1-persistente e configurações harmônicas, apresenta-se como uma estratégia descentralizada robusta. O modelo otimiza as métricas de AoI e mitiga os efeitos de colisões em larga escala, cumprindo os requisitos de conectividade sem depender de gerenciamento centralizado.

⁵O conteúdo desta seção consiste em um resumo adaptado pelo pesquisador *Plínio S. Dester*, com base em dois artigos publicados no *IEEE Communications Letters*, intitulados “Analysis of the 1-Persistent Age Threshold Slotted ALOHA” e “On the Deterministic Regime of 1-Persistent Age-Threshold Slotted ALOHA in Heterogeneous Networks” [67, 68].

7.2 Modelo do Sistema

Topologia da Rede e Modelo de Tráfego

Consideramos uma rede de comunicação centralizada composta por um conjunto de $n \in \mathbb{Z}_{\geq 2}$ nós (dispositivos IoT/sensores) que transmitem atualizações de status para uma única estação base (*Base Station* (BS)) em um canal compartilhado. O tempo é discretizado em slots de duração unitária, indexados por $t \in \mathbb{Z}_{\geq 0}$.

Assume-se que os dispositivos operam sob o modelo de tráfego *generate-at-will* (ou geração sob demanda). Neste modelo, evita-se a necessidade de gerenciar filas locais de pacotes, pois cada nó amostra e gera uma nova atualização imediatamente antes de decidir transmiti-la. Consequentemente, as atualizações transmitidas contêm sempre a informação mais recente disponível no nó.

AoI e Evolução de Estado

A *idade* das informações recebidas pela estação base é representada pela métrica AoI. Seja $U_i(t)$ o instante de geração do pacote mais recentemente recebido (até o tempo t) proveniente do nó i . A AoI do nó i no tempo t é definida como

$$\Delta_i(t) \triangleq t - U_i(t), \quad (35)$$

sendo o símbolo $\triangleq$ “igual por definição”.

Sejam os processos estocásticos de tempo discreto $e_i(t), \theta_i(t) \in \{0, 1\}$ indicando, respectivamente, se o nó i acessa o canal no slot t e se houve transmissão bem-sucedida do nó i . Sob tráfego *generate-at-will*, a evolução do processo da AoI para cada nó i segue a dinâmica

$$\Delta_i(t+1) = \begin{cases} 1, & \text{se } \theta_i(t) = 1; \\ \Delta_i(t) + 1, & \text{se } \theta_i(t) = 0. \end{cases} \quad (36)$$

A partir dessa evolução, definem-se as métricas de desempenho no estado estacionário para cada nó, como o throughput η_i , a AoI média no tempo $\overline{\text{AoI}}_i$ e a PAoI média no tempo $\overline{\text{PAoI}}_i$ para o i -ésimo nó, respectivamente, por

$$\eta_i \triangleq \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \theta_i(t), \quad (37)$$

$$\overline{\text{AoI}}_i \triangleq \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \Delta_i(t), \quad (38)$$

$$\overline{\text{PAoI}}_i \triangleq \lim_{T \rightarrow \infty} \frac{\sum_{t=0}^{T-1} \Delta_i(t) \theta_i(t)}{\sum_{t=0}^{T-1} \theta_i(t)}, \quad (39)$$

se os limites existirem. Ademais, o throughput da rede é definido como

$$\eta \triangleq \sum_{i=1}^n \eta_i. \quad (40)$$

Modelo de Acesso

Para gerenciar o acesso ao meio, a rede emprega uma variante descentralizada do protocolo TSA, denominada 1-persistent TSA. Neste esquema, cada nó i possui um limiar de idade pré-determinado, denotado por $A_i \in \mathbb{Z}_{>0}$, que define seu requisito específico de AoI.

Portanto, a variável de acesso $e_i(t)$ é expressa matematicamente por

$$e_i(t) = \mathbb{1}_{\{\Delta_i(t)=A_i\}} + B_i(t)\mathbb{1}_{\{\Delta_i(t)>A_i\}}, \quad (41)$$

sendo $\mathbb{1}_{\{\cdot\}}$ a função indicadora e $B_i(t)$ um processo de Bernoulli com parâmetro $p \in (0,1)$, independente entre os nós. Esta regra baseia-se exclusivamente no valor local de $\Delta_i(t)$ e reflete o seguinte comportamento:

1. **Silêncio:** Se $\Delta_i(t) < A_i$, a informação ainda é considerada “fresca” o suficiente e o nó i não transmite ($e_i(t) = 0$).
2. **Acesso 1-persistente:** No exato momento em que a idade atinge o limiar ($\Delta_i(t) = A_i$), o nó i transmite uma nova atualização com probabilidade 1.
3. **Retransmissão:** Se a transmissão inicial falhar e a idade ultrapassar o limiar ($\Delta_i(t) > A_i$), o nó i entra em um estado de retransmissão, tentando acessar o canal com probabilidade p em cada slot subsequente.

O modelo abrange tanto redes homogêneas (em que $A_i = A$ para todo i) quanto redes heterogêneas (em que os nós podem ser configurados com conjuntos distintos de limiares, $A_1, \dots, A_n$, para refletir diferentes requisitos de frescor da informação).

Modelo de Canal e Colisões

A camada física é modelada como um canal de colisão tradicional (*collision channel*). Define-se que uma transmissão realizada pelo nó i no slot t é decodificada com sucesso pela BS se, e somente se, nenhum outro nó $j \neq i$ transmita no mesmo slot. Portanto, a variável de sucesso da transmissão $\theta_i(t)$ é dada por

$$\theta_i(t) = \mathbb{1}_{\{\sum_{k \neq i} e_k(t)=0, e_i(t)=1\}}. \quad (42)$$

Se dois ou mais nós decidirem acessar o canal concorrentemente, ocorre uma colisão, resultando na perda de todas as atualizações envolvidas. Neste caso de falha, os nós em colisão não reiniciam suas idades e passam a competir pelo canal no regime de retransmissão nos slots seguintes.

Vale destacar que, nas simulações, considera-se um canal de apagamento em que uma transmissão sofre apagamento com probabilidade ϵ (independentemente dos demais timeslots), mesmo na ausência de colisão.

7.3 Resultados Limites

Nesta seção, caracteriza-se o regime estacionário da cadeia de Markov associada ao protocolo 1-persistent TSA em redes heterogêneas. Seja a cadeia de Markov com espaço de estados $\mathcal{S} = \mathbb{Z}_{\geq 0}^n$, representando as idades dos nós, e o processo estocástico no tempo discreto $t \in \mathbb{Z}_{\geq 0}$ definido por

$$S(t) \triangleq (\Delta_1(t), \dots, \Delta_n(t)) \in \mathcal{S}. \quad (43)$$

A análise da convergência da rede para um regime estacionário ordenado e livre de colisões revela uma conexão geométrica e temporal com os limiares $A_1, A_2, \dots, A_n$ formarem uma sequência harmônica, cuja definição é apresentada a seguir.

Sequências Harmônicas e Conjunto Invariante

Definição 1 (Sequência Harmônica). Uma sequência finita $A_1, A_2, \dots, A_n$ de inteiros positivos é chamada de sequência harmônica se existirem deslocamentos (*offsets*) $x_1, x_2, \dots, x_n \in \mathbb{Z}$ tais que as progressões aritméticas $\mathcal{A}_i \triangleq \{x_i + A_i k \mid k \in \mathbb{Z}\}$ sejam disjuntas aos pares, isto é, $\mathcal{A}_i \cap \mathcal{A}_j = \emptyset$ para todo $i \neq j$.

Em outras palavras, uma sequência é harmônica⁶ se as progressões aritméticas puderem ser alinhadas de modo que suas repetições nunca coincidam, formando uma coleção de padrões periódicos não interferentes sobre os números inteiros. Por exemplo, a sequência 2, 3 não é harmônica, mas as sequências 2, 2 e 2, 4, 4 são harmônicas.

Lema 3 ([70]). Se $A_1, \dots, A_n$ é uma sequência harmônica, então a carga dos limiares satisfaz

$$\rho \triangleq \sum_{i=1}^n \frac{1}{A_i} \leq 1 \quad e \quad \gcd(A_i, A_j) > 1, \quad \forall i \neq j.$$

O operador $\gcd(a, b)$ denota o máximo divisor comum entre a e b . A quantidade ρ representa a proporção de inteiros cobertos pelo conjunto $\cup_i \mathcal{A}_i$ e, evidentemente, não pode exceder 1. No contexto desta seção, ρ também corresponde à fração de intervalos de tempo alocados pelas transmissões nos limiares de idade e é denominada *aggregate threshold load*.

Uma condição suficiente elegante para que uma sequência seja harmônica é que os máximos divisores comuns sejam distintos dois a dois e diferentes de 1 [70]. Caracterizar todas as sequências harmônicas, ou seja, fornecer uma condição necessária e suficiente completa para que uma sequência de inteiros positivos seja harmônica, permanece um problema em aberto da Teoria dos Números [70].

A partir da Definição 1, estabelece-se o subconjunto $\mathcal{S}_A \subset \mathcal{S}$, que possui os estados nos quais os processos de idade estão alinhados de modo a evitar colisões nos limiares:

$$\mathcal{S}_A \triangleq \{(x_1, \dots, x_n) \in \mathcal{S} \mid x_i + A_i q \neq x_j + A_j r, \quad \forall i \neq j, \quad \forall q, r \in \mathbb{Z}, \quad 1 \leq x_k \leq A_k \quad \forall k\}.$$

Pela Definição 1, o conjunto $\mathcal{S}_A$ é não-vazio se, e somente se, a sequência de limiares for harmônica. Para descrever a dinâmica livre de colisões no interior de $\mathcal{S}_A$, utiliza-se o mapeamento determinístico para o estado sucessor, $f_A : \mathcal{S}_A \rightarrow \mathcal{S}_A$, definido como

$$f_A((x_i)_{i=1}^n) = ((x_i \bmod A_i) + 1)_{i=1}^n,$$

sendo o operador $a \bmod b$ o resto da divisão de a por b .

Lema 4. Se $S(t) \in \mathcal{S}_A$ para algum $t \in \mathbb{Z}_{\geq 0}$, então $S(t+1) = f_A(S(t)) \in \mathcal{S}_A$.

Demonstração. Seja $S(t) = x = (x_1, \dots, x_n) \in \mathcal{S}_A$. Pela definição de $\mathcal{S}_A$, temos $1 \leq x_i \leq A_i$ para todo i , e para todo $i \neq j$,

$$x_i + A_i q \neq x_j + A_j r, \quad \forall q, r \in \mathbb{Z}. \quad (44)$$

⁶Sequência harmônica não deve ser confundida com progressão harmônica.

Sob a regra do 1-persistent TSA, o nó i transmite com probabilidade um se, e somente se, $\Delta_i(t) = A_i$. Logo, se $x_i < A_i$, o nó i permanece em silêncio e $\Delta_i(t+1) = x_i + 1$; se $x_i = A_i$, o nó i transmite. Além disso, dois nós distintos não podem satisfazer simultaneamente $x_i = A_i$ e $x_j = A_j$, pois isso violaria (44) com $q = r = 0$. Portanto, a transmissão é livre de colisões e bem-sucedida, resultando em $\Delta_i(t+1) = 1$. Assim, para todo i , $\Delta_i(t+1) = (x_i \bmod A_i) + 1$, e consequentemente $S(t+1) = f_A(S(t))$.

Seja $y = f_A(x)$, ou seja, $y_i = (x_i \bmod A_i) + 1$. Claramente, $1 \leq y_i \leq A_i$ para todo i . Resta mostrar que a propriedade de invariância se mantém, isto é, para todo $i \neq j$, $y_i + A_i q \neq y_j + A_j r$ para quaisquer $q, r \in \mathbb{Z}$. Defina $\varepsilon_i = 1$ se $x_i = A_i$, e $\varepsilon_i = 0$ caso contrário. Então, $y_i + A_i q = x_i + 1 + A_i(q - \varepsilon_i)$. Se a igualdade $y_i + A_i q = y_j + A_j r$ fosse válida, implicaria que $x_i + A_i(q - \varepsilon_i) = x_j + A_j(r - \varepsilon_j)$, o que contradiz (44), já que $q - \varepsilon_i$ e $r - \varepsilon_j$ são inteiros. Portanto, tal igualdade é impossível e o estado sucessor y pertence de fato a $\mathcal{S}_A$. $\square$

Uma vez em $\mathcal{S}_A$, o sistema torna-se periódico, como mostrado no lema a seguir.

Lema 5. *Se $S(t_0) \in \mathcal{S}_A$ para algum $t_0 \in \mathbb{Z}_{\geq 0}$, então $S(t) \in \mathcal{S}_A$ e a evolução temporal das idades é periódica, obedecendo $\Delta_i(t) = \Delta_i(t + A_i)$ para todo $t \geq t_0$ e $i \in \{1, \dots, n\}$.*

Demonstração. Aplicando o Lema 4 recursivamente, obtemos $S(t) = f_A^{t-t_0}(S(t_0)) \in \mathcal{S}_A$ para todo $t \geq t_0$, provando a primeira afirmação. Considere agora a evolução da idade do nó i . Como a dinâmica dentro de $\mathcal{S}_A$ é estritamente determinística e governada pelo mapeamento f_A , a idade do nó i aumenta em uma unidade a cada passo de tempo até atingir A_i , ponto em que é redefinida para 1, repetindo esse padrão periodicamente. Logo, para todo $t \geq t_0$, $\Delta_i(t + A_i) = \Delta_i(t)$, provando que o processo de idade de cada nó i é periódico com período de A_i slots. $\square$

Convergência e Desempenho Estacionário

O teorema a seguir fundamenta a principal garantia operacional da rede: caso os limiares de idade formem uma sequência harmônica, o sistema abandona a fase de contenção (mesmo iniciando em um estado congestionado ou após uma perturbação no canal) e converge *almost surely* (a.s.) para um comportamento determinístico e periódico (análogo ao TDMA).

Teorema 3. *Seja $S(0) \in \mathcal{S}$. Temos que, a.s.,*

$$\lim_{t \rightarrow \infty} \mathbb{1}_{\{S(t) \in \mathcal{S}_A\}} = 1$$

se, e somente se, a sequência $A_1, A_2, \dots, A_n$ for harmônica.

Demonstração. Considere uma cadeia de Markov agregada com espaço de estados

$$\mathcal{S}' \triangleq \{(x_1, \dots, x_n) \in \mathcal{S} \mid x_i \leq A_i + 1 \ \forall i\},$$

que satura as idades nos sucessores dos limiares. A agregação é válida pois todos os estados onde a idade ultrapassa o limiar ($\Delta_i > A_i$) induzem probabilidades de transição idênticas devido à natureza do processo de retransmissão de Bernoulli. Como o espaço de estados $\mathcal{S}'$ é finito e, para todo estado inicial $S(0) \in \mathcal{S}' \setminus \mathcal{S}_A$, existe um caminho com probabilidade estritamente positiva que leva ao conjunto invariante $\mathcal{S}_A$ (uma sequência bem escolhida de retransmissões é suficiente), o resultado segue imediatamente do Lema 5. O caso em que $S(0) \in \mathcal{S}_A$ é trivial. A recíproca também é trivial: se $A_1, \dots, A_n$ não forma uma sequência harmônica, a Definição 1 impõe que $\mathcal{S}_A$ seja um conjunto vazio, inviabilizando a convergência. $\square$

Nota 1. No caso homogêneo em que $A_1 = A_2 = \dots = A_n$ e $\rho \leq 1$, os limiares são sempre harmônicos e, pelo Teorema 3, a convergência para um regime determinístico ocorre a.s.

Corolário 3. Se $A_1, \dots, A_n$ for uma sequência harmônica, as métricas de desempenho limite da rede tornam-se independentes da probabilidade de retransmissão p_i e resultam, a.s., em

$$\eta_i = \frac{1}{A_i}, \quad \overline{\text{AoI}}_i = \frac{A_i + 1}{2}, \quad \overline{\text{PAoI}}_i = A_i.$$

Ademais, a vazão global da rede atinge $\eta = \rho = \sum_{i=1}^n \frac{1}{A_i} \leq 1$ (a.s.).

Demonstração. A demonstração consiste em combinar o Teorema 3 com os caminhos periódicos do Lema 5. Uma vez no regime estacionário contido em $\mathcal{S}_A$, o nó i acessa o canal perfeitamente 1 vez a cada A_i slots. Aplicando as definições das métricas sobre essa periodicidade ideal, as expressões analíticas das métricas convergem a.s. para as expressões mostradas. $\square$

Nota 2. Do Corolário 3, vale notar que as métricas de performance não dependem da probabilidade de retransmissão p . Por outro lado, p influencia o tempo necessário para o sistema convergir para $\mathcal{S}_A$. De fato, há um intervalo apropriado para p que equilibra a taxa de tentativas com o número de colisões, minimizando a duração da fase transitória.

Nota 3. Utilizando o fato de que a média aritmética é maior ou igual à média harmônica no Corolário 3, obtemos os seguintes limitantes inferiores para o AoI e PAoI médios globais:

$$\sum_{i=1}^n \frac{\overline{\text{AoI}}_i}{n} \geq \frac{n + \rho}{2\rho}, \quad \sum_{i=1}^n \frac{\overline{\text{PAoI}}_i}{n} \geq \frac{n}{\rho}. \quad (45)$$

Com igualdade apenas no caso homogêneo ($A_1 = A_2 = \dots = A_n$).

Construção de Limiares Harmônicos

Para viabilizar o projeto e a alocação prática de limiares harmônicos heterogêneos, apresenta-se, a seguir, uma regra construtiva simples.

Teorema 4. Sejam $b, c \in \mathbb{Z}_{>0}$. Seja $A_1, A_2, \dots, A_n$ uma sequência finita de inteiros positivos na forma

$$A_i = c b^{k_i}, \quad \text{para } k_1, \dots, k_n \in \mathbb{Z}_{\geq 0}.$$

A referida sequência $A_1, A_2, \dots, A_n$ é harmônica se, e somente se, sua carga agregada $\rho \leq 1$.

Demonstração. A necessidade da condição decorre diretamente da propriedade global das sequências harmônicas (a carga não pode exceder 1). Provaremos a suficiência de forma construtiva. Suponha que $\rho \leq 1$, ou seja, $\sum_i b^{-k_i} \leq c$. Vamos derivar inteiros x_i de modo que as progressões sejam disjuntas aos pares. Ordene os índices $\{i\}_{i=1}^n$ em ordem decrescente de b^{-k_i} e distribua-os de forma gulosa em uma partição de c subconjuntos $G_0, \dots, G_{c-1}$ tais que cada partição satisfaça a restrição $\sum_{i \in G_r} b^{-k_i} \leq 1$. Para cada subconjunto G_r , existe uma sequência de inteiros $(a_i)_{i \in G_r}$ tal que os conjuntos $I_i = [\frac{a_i}{b^{k_i}}, \frac{a_i+1}{b^{k_i}})$, para $i \in G_r$, formem uma coleção de subintervalos não-sobrepostos em $[0, 1)$. Definindo os deslocamentos como $x_i = r + ca_i$, a progressão resultante torna-se $\mathcal{A}_i = \{r + c(a_i + b^{k_i}q) \mid q \in \mathbb{Z}\}$. Se $r \neq s$, o resto da divisão por c difere e garante a disjunção dos conjuntos associados a G_r e G_s . Se $r = s$ e $k_i \leq k_j$, a disjunção de I_i e I_j , aliada ao fato de que b^{k_i} divide b^{k_j} , implica que $(a_j - a_i) \bmod b^{k_i} \neq 0$, e, consequentemente, $(x_j - x_i) \bmod (cb^{k_i}) \neq 0$. Logo, todas as progressões são disjuntas, o que comprova que a sequência construída é harmônica. $\square$

Nota 4. Usando o Teorema 4, uma estratégia de projeto simples para construir uma sequência harmônica de limiares consiste em fixar uma base (por exemplo, $b = 2$) e atribuir os expoentes $k_1, k_2, \dots, k_n$ de acordo com os níveis de urgência dos nós: os mais críticos recebem o expoente 0, os do próximo nível, o expoente 1, e assim por diante. O parâmetro de escala c é então escolhido de modo que, em condições nominais, a rede opere sob uma carga desejada ρ .

7.4 Análise do Transiente

A análise do comportamento transiente em redes heterogêneas sob o protocolo 1-persistent TSA apresenta uma complexidade analítica extrema, devido à intrincada relação entre os estados de cada nó. Por essa razão, nesta seção, assume-se uma rede homogênea ($A_i = A$ para todo i). Mesmo nesse cenário simplificado, a modelagem exata do transiente permanece desafiadora, o que motiva o uso de um modelo aproximado baseado na evolução da probabilidade de colisão.

Nesta seção, define-se uma cadeia de Markov auxiliar de tempo discreto $\{K(t)\}_{t \in \mathbb{Z}_+}$ com espaço de estados $\mathcal{K} \triangleq \{0, 1, \dots, n\}$. O estado $K(t) = k$ representa um estado abstrato do sistema indexado por k , e a cadeia é definida pelas seguintes probabilidades de transição. Para $k \in \mathcal{K}$,

$$p_k^+ = \frac{n-k}{A}(1 - (1-p)^k), \quad (46)$$

$$p_k^- = \left(1 - \frac{n-k}{A}\right) kp(1-p)^{k-1}, \quad (k \neq 0) \quad (47)$$

$$p_k^o = 1 - p_k^+ - p_k^-, \quad (48)$$

onde p_k^+ , p_k^- , e p_k^o denotam as probabilidades de transição do estado k para os estados $k+1$, $k-1$, e k , respectivamente.

Esta cadeia auxiliar é um modelo simplificado, construído para aproximar o comportamento transiente do processo original $\{S(t)\}_t$ quando os usuários possuem idades acima do limiar. Em particular, a absorção de $\{K(t)\}_t$ no estado 0 corresponde ao sistema atingir uma configuração na qual nenhum usuário permanece acima do limiar, o que é o evento de interesse na análise do transiente e está associado ao processo original atingir $\mathcal{S}_A$.

Define-se que o regime transiente termina no instante em que todos os nós têm uma idade menor ou igual a A , o que ocorre após atingir o conjunto $\mathcal{S}_A$ no modelo original de Markov $\{S(t)\}$, e o estado 0 no modelo aproximado de Markov $\{K(t)\}$. A seguinte proposição fornece um resultado analítico para calcular uma aproximação da duração esperada do transiente começando com k usuários com idades maiores que A .

Proposição 1. *Considere a cadeia de Markov $\{K(t)\}$. Seja o tempo esperado para atingir o único estado absorvente $0 \in \mathcal{K}$ começando do estado $k \in \mathcal{K}$ denotado por E_k . Então,*

$$E_k = \sum_{i=1}^k D_i, \quad k \in \mathcal{K} \setminus \{0\}, \quad E_0 = 0,$$

onde $\{D_k\}_{k \in \{1, 2, \dots, n\}}$ é a única solução da seguinte equação de diferenças, para $k \in \{1, 2, \dots, n\}$,

$$\left(1 - \frac{n-k}{A}\right) kp(1-p)^{k-1} D_k - \frac{n-k}{A}(1 - (1-p)^k) D_{k+1} = 1.$$

Demonstração. Para $k \in \{1, \dots, n-1\}$,

$$E_k = 1 + p_k^+ E_{k+1} + p_k^- E_{k-1} + p_k^0 E_k,$$

o que pode ser rearranjado como

$$p_k^- (E_k - E_{k-1}) - p_k^+ (E_{k+1} - E_k) = 1.$$

Notar que $D_k = E_k - E_{k-1}$ conclui a primeira parte da prova. A unicidade segue do fato de que a equação de diferenças possui uma solução única. Avaliando-a em $k = n$, temos

$$D_n = \frac{1}{np(1-p)^{n-1}}. \quad (49)$$

Então, é direto determinar D_{k-1} a partir de D_k para $k = n, n-1, \dots, 2$. Isso prova a unicidade. $\square$

Corolário 4. *Considere a cadeia de Markov $\{K(t)\}$. Se o parâmetro A aumentar, então o tempo esperado E_k diminui para todo $k \in \mathcal{K} \setminus \{0\}$.*

Demonstração. Note a partir de (49) que D_n não muda com A . Então, é fácil mostrar a partir da equação de recorrência da Proposição 1 que D_{n-1} diminui com A . Analogamente, é fácil provar que D_{k-1} diminui com A se D_k diminui com A . Evocar o princípio da indução matemática completa a prova porque $E_k = \sum_i D_i$. $\square$

Assim, há um *trade-off* na seleção do limiar de idade ótimo $A^* \geq n$, pois limiares menores melhoram o desempenho, enquanto limiares maiores encurtam o período transiente. Uma abordagem simples é selecionar o menor valor para o qual a probabilidade de transição do estado 1 para 0 é maior que a probabilidade de transição de 1 para 2. Após algumas manipulações algébricas da desigualdade $p_1^+ < p_1^-$, temos

$$A^* = 2n - 1, \quad (50)$$

que, convenientemente, não depende de p .

Como a probabilidade de acesso p não influencia as métricas de desempenho no 1-persistent TSA, a escolha da probabilidade de acesso ótima p^* deve focar primariamente na minimização da duração do transiente. Propomos a otimização:

$$p^* = \arg \min_p E_n, \quad (51)$$

que minimiza a duração esperada do transiente a partir do estado de pior caso $k = n$. Usando evidências numéricas com n variando de 50 a 500, encontramos que $p \approx 2,5/n$ é próximo do ótimo quando o limiar de idade $A = 2n - 1$.

Nota 5. Considere um sistema com n nós operando em um estado em $\mathcal{S}_A$. Se m nós adicionais entrarem no sistema, resultando em $n' = n + m \leq A$ nós, o tempo esperado para alcançar $\mathcal{S}_A$ novamente é dado por E_m avaliado para um sistema de tamanho n' . Por outro lado, se m nós deixarem o sistema, o estado do sistema permanece em $\mathcal{S}_A$.

7.5 Simulações

Rede Homogênea

Na Figura 18, são apresentados os valores teóricos da duração média do transitório E_n de acordo com o modelo aproximado $\{K(t)\}$, partindo do estado $k = n$, e comparados com os resultados de simulação do modelo original $\{S(t)\}$, iniciando com todas as idades iguais a $A + 1$. Constata-se que a formulação aproximada captura a tendência da duração transitória da cadeia de Markov original.

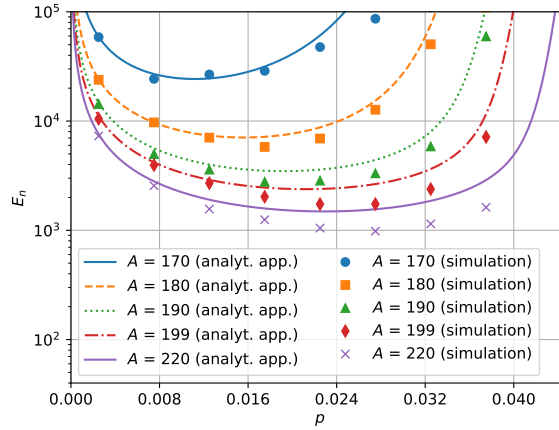
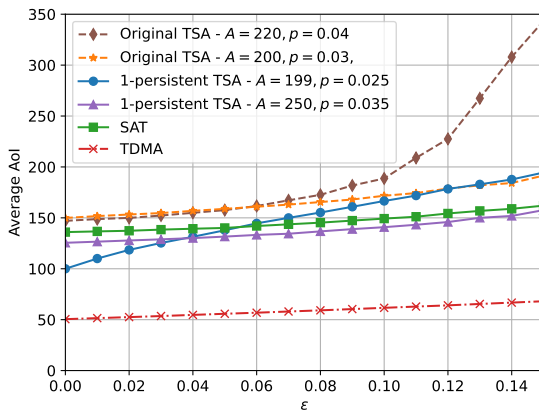
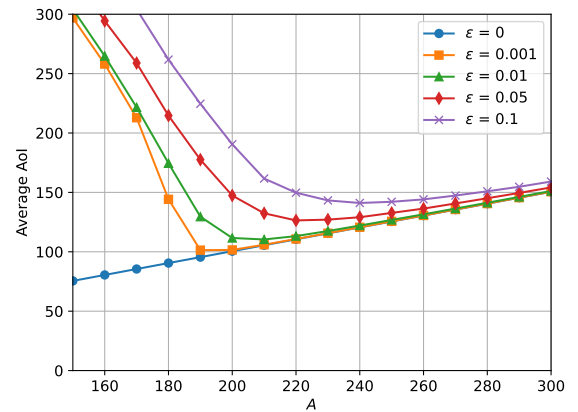


Figura 18: Duração média do transiente E_n em função da probabilidade de acesso p para diversos valores do limiar A e $n = 100$ nós. Os pontos representam os resultados da simulação do modelo original $\{S(t)\}$, e as curvas os resultados analíticos do modelo aproximado $\{K(t)\}$.



(a) AoI médio em função de ε .



(b) AoI médio em função de A .

Figura 19: AoI médio do 1-persistent TSA para $n = 100$ nós.

Considerando um canal de apagamento em que uma transmissão pode falhar com probabilidade $\varepsilon > 0$ mesmo na ausência de colisões, o Teorema 3 deixa de ser válido, e a convergência para um comportamento do tipo TDMA deixa de ser garantida. Ainda assim, os resultados de simulação da Figura 19a mostram que o 1-persistent TSA supera um importante *benchmark*

para protocolos baseados em AoI denominado *Stationary Age-based Thinning* (SAT) [71] e, conseqüentemente, o TSA original, na faixa de probabilidades de erro considerada.

Além disso, a robustez a erros de transmissão melhora para valores maiores do limiar de idade A . Como ilustrado na Figura 19b, o limiar que minimiza a AoI desloca-se para valores maiores à medida que ε aumenta, uma vez que limiares maiores reduzem as durações transitórias (Corolário 4) e os erros de transmissão atuam como perturbações ocasionais. Embora os resultados teóricos sejam válidos para $A \geq n = 100$, as simulações indicam que, para $A \leq 180$, o esquema torna-se altamente sensível a pequenas probabilidades de erro, evidenciando a necessidade de uma escolha equilibrada de A que estabeleça um compromisso entre o desempenho em regime permanente, a duração transitória e a robustez.

Rede Heterogênea

Como o 1-persistent TSA supera o TSA original e o protocolo SAT em redes homogêneas, o caso homogêneo é adotado como *benchmark*. Define-se também a *PAoI média normalizada*

$$\overline{\text{NPAoI}} \triangleq \frac{1}{n} \sum_{i=1}^n \frac{\overline{\text{PAoI}}_i}{A_i}, \quad (52)$$

que representa a PAoI média de cada nó normalizada pelo seu requisito de AoI.

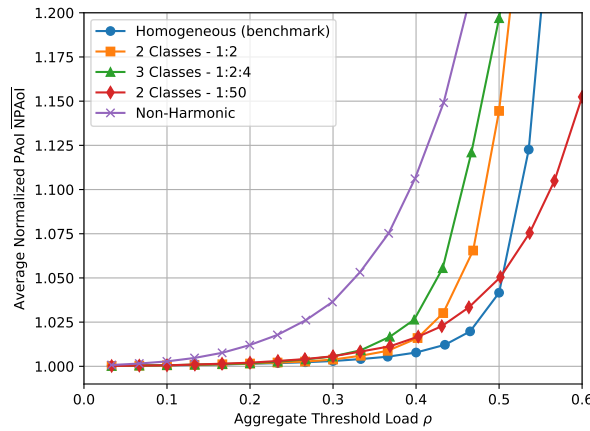


Figura 20: AoI médio em função de ρ com $n = 60$ nós e $p = 2,5/n$.

Como o limiar A_i reflete o ponto de operação desejado de idade para o nó i , A_i serve naturalmente como referência para a normalização. Nas simulações, para avaliar a robustez do protocolo, considera-se uma probabilidade de apagamento $\varepsilon = 1\%$ para modelar perturbações ocasionais no canal. A Figura 20 apresenta os resultados de simulação. No cenário *3 Classes – 1:2:4*, os nós são divididos em três subconjuntos de mesmo tamanho. O primeiro, o segundo e o terceiro subconjuntos utilizam limiares c , $2c$ e $4c$, respectivamente ($b = 2, k_i = 0, 1, 2$); o parâmetro c é variado para cobrir a faixa de ρ investigada. As configurações *2 Classes – 1:2* e *2 Classes – 1:50* são definidas de maneira análoga. Pelo Teorema 4, todas essas configurações são harmônicas. Em contraste, o cenário *Non-Harmonic* também divide os nós em dois grupos de mesmo tamanho, mas o segundo limiar excede o primeiro em uma unidade. Pelo Lema 3, essa configuração não é harmônica, apesar de ser próxima do caso homogêneo. Pelo Corolário 3, na ausência de perturbações e sob configurações harmônicas de limiares, a PAoI média normalizada

satisfaz $\overline{\text{NPAoI}} = 1$ para todo $\rho \leq 1$. Para valores moderados de ρ , os resultados de simulação corroboram esse comportamento teórico. Entretanto, à medida que ρ aumenta, perturbações ocasionais já são suficientes para impedir que o sistema opere próximo ao regime estacionário do tipo TDMA. Esse comportamento é esperado, pois o tempo necessário para o sistema atingir $\mathcal{S}_A$ cresce à medida que sua cardinalidade $|\mathcal{S}_A|$ diminui, tornando o protocolo mais sensível a perturbações.

Como os cenários considerados envolvem diferentes conjuntos de requisitos de AoI, as comparações são delicadas. Porém, duas conclusões podem ser extraídas. Primeiro, a violação da condição de harmonicidade leva a uma degradação acentuada de desempenho: mesmo quando os limiares são próximos aos da configuração homogênea, o desempenho resultante é inferior ao de todos os casos harmônicos. Segundo, aumentar a separação entre limiares de idade harmônicos reduz a robustez do sistema, dificultando o atendimento aos requisitos na presença de perturbações e de tráfego elevado.

7.6 Conclusão

Nesta atividade, apresentou-se o desenvolvimento e a análise do protocolo 1-persistent TSA, proposto como uma solução descentralizada e eficiente para a concepção da rede de acesso aleatório em cenários de conectividade massiva (MTC e IoT) para redes 6G. Ao introduzir uma modificação simples, porém inovadora, que força a transmissão com probabilidade unitária no instante exato em que o limiar de idade é atingido, o protocolo busca capturar a eficiência de um esquema determinístico sem os elevados custos de sinalização e coordenação centralizada.

A caracterização teórica do estado estacionário revelou propriedades fundamentais do protocolo em redes heterogêneas. Demonstrou-se, matematicamente, que o sistema converge *almost surely* para um regime determinístico e livre de colisões, análogo ao TDMA, se, e somente se, o conjunto de limiares de idade configurar uma sequência harmônica. Sob essa condição de alinhamento temporal, as métricas de throughput, AoI média e PAoI média assumem expressões em forma fechada simples e tornam-se estritamente independentes da probabilidade de retransmissão. Por outro lado, as investigações confirmaram que a violação da condição de harmonicidade compromete gravemente o desempenho do sistema devido ao surgimento de colisões sistemáticas.

Adicionalmente, a análise do comportamento transiente explicitou um *trade-off* de projeto crucial para redes homogêneas: a minimização da AoI em regime permanente (via redução do limiar A) acarreta um prolongamento severo da fase transitória de contenção. O modelo aproximado em cadeia de Markov proposto mostrou-se uma ferramenta eficaz para mapear essa dinâmica, permitindo derivar diretrizes práticas, como o limiar de compromisso $A^* = 2n - 1$ e a probabilidade de retransmissão quase-ótima $p^* \approx 2,5/n$.

Por fim, os resultados de simulação validaram a robustez do 1-persistent TSA mesmo em canais ruidosos sujeitos a apagamentos. O protocolo superou consistentemente *benchmarks* relevantes da literatura, como o protocolo SAT e o TSA original, mostrando-se uma estratégia robusta e promissora para otimizar o frescor da informação e assegurar a eficiência de redes de comunicação de próxima geração.

8 Fake Base Station: Ameaças, Detecção e Mitigação.

Christian Mailer, Stevan Grubisic, Adão Boava, Hirley Alves, Richard Demo Souza.
UFSC

8.1 Introdução

As redes de comunicações móveis evoluíram significativamente desde a Segunda Geração de Rede Móvel Celular (2G), passando a suportar não apenas serviços de voz e mensagens curtas, mas também aplicações críticas como controle industrial, veículos conectados e comunicações de missão crítica. Com a Quinta Geração de Rede Móvel Celular (5G), surgiram novas classes de serviço, cada uma com requisitos rigorosos de confiabilidade, latência e disponibilidade.⁷

Apesar dessas evoluções, ataques baseados em *Fake Base Stations* (FBSs), também conhecidas como *rogue* ou *false base stations*, continuam viáveis. Uma FBS é uma estação rádio base operando fora do controle da operadora, capaz de explorar estágios iniciais de acesso à rede e outras vulnerabilidades para induzir comportamentos maliciosos em um *User Equipment* (UE). Na prática, uma FBS pode ser construída com hardware de custo relativamente baixo, como um *Software Defined Radio* (SDR) e um computador portátil, combinados com pilhas de *software open-source* para a rede de acesso e o núcleo. Todo o conjunto pode ser transportado em uma mochila ou veículo, permitindo ataques móveis que atingem grande número de usuários com baixo risco de detecção. Equipamentos comerciais de FBS, originalmente destinados a agências de inteligência e forças de segurança, também estão disponíveis no mercado, embora alguns fornecedores os comercializem sem restrição de uso.

A persistência desses ataques é agravada pela convivência multigeracional das redes. A relevância prática da ameaça é evidenciada pelo cenário brasileiro recente: ataques baseados em FBS para envio de *Short Message Service* (SMS) fraudulentos, antes restritos a países como a China, tornaram-se frequentes em grandes centros urbanos do Brasil, motivando ações conjuntas entre a Anatel e forças de segurança pública para localizar e desativar equipamentos clandestinos. A estratégia típica dos atacantes envolve duas etapas: primeiro, uma FBS transmitindo na Quarta Geração de Rede Móvel Celular (4G) ou 5G força o UE a realizar um *downgrade* para 2G; em seguida, outra FBS transmitindo em 2G executa o envio de mensagens falsificadas, explorando a ausência de autenticação mútua nessa geração.

Do ponto de vista do impacto, ataques baseados em FBS não se limitam a fraudes financeiras. Ataques de disrupção podem afetar comunicações de emergência, sistemas de alerta público, veículos autônomos e dispositivos industriais, elevando o risco de danos físicos. À medida que o 6G busca requisitos ainda mais exigentes, compreender e resolver as vulnerabilidades herdadas de gerações anteriores torna-se essencial para que as arquiteturas de segurança de próxima geração não perpetuem fragilidades já conhecidas.

Este capítulo apresenta uma visão geral do conteúdo de um artigo do tipo *survey*, em preparação, que adota uma abordagem FBS-cêntrica, organizando a análise a partir dos métodos de implantação das FBS e dos tipos de ataque que estes habilitam, com foco nas vulnerabilidades exploráveis. O presente capítulo resume os aspectos centrais desse trabalho; detalhes adicionais poderão ser consultados no artigo quando este estiver finalizado e disponibilizado online.

⁷O conteúdo desta seção consiste em um resumo adaptado pelo pesquisador Christian Mailer, com base em um artigo do tipo *survey*, ainda em preparação, intitulado: “Fake Base Station Threats from 2G to 5G: Attacks, Defenses, and Lessons for 6G”.

8.2 Métodos de Ataque Baseados em FBS

Métodos de ataque descrevem *como* uma FBS é implantada e operada no ambiente celular. Um mesmo método pode ser utilizado para alcançar diferentes objetivos maliciosos.

Man-in-The-Middle (MiTM): O adversário implanta uma FBS em conjunto com um *Fake User Equipment* (FUE), criando um *relay* transparente entre o UE legítimo e a estação rádio base legítima. Todo o tráfego passa pela infraestrutura controlada pelo atacante, que pode observar, atrasar, descartar ou modificar mensagens, especialmente antes da ativação de mecanismos de integridade. O *relay* pode ser co-localizado, minimizando latência, ou distribuído, com os dois componentes interconectados por um enlace de alta velocidade para ampliar o alcance. Em cenários de ondas milimétricas, onde a cobertura é limitada a poucas centenas de metros, a implantação torna-se mais simples, pois o atacante pode posicionar o *relay* entre o UE e uma estação base fora do seu alcance direto. O principal desafio é manter latência suficientemente baixa para evitar expiração de temporizadores de protocolo.

Injeção de Mensagens: Consiste na transmissão pontual de mensagens falsas, sincronizadas com a comunicação legítima entre UE e rede. Diferentemente do MiTM, não exige conexão persistente, sendo energeticamente eficiente e difícil de detectar, pois o injetor não transmite sinais de sincronização nem mensagens *broadcast*. Mensagens *Non-Access Stratum* (NAS) ou RRC falsas podem induzir rejeições, *downgrades* ou falhas de registro. O atacante deve, contudo, estar sincronizado com a estação base e o UE para injetar no instante correto.

Personificação: A FBS se faz passar por uma estação base legítima, ou o FUE se faz passar por um UE legítimo perante a rede. A personificação é facilitada quando a rede ou o UE aceita algoritmos nulos de *ciphering* ou integridade fora de contextos de emergência, contrariando recomendações do *3rd Generation Partnership Project* (3GPP). Avaliações empíricas revelaram que algumas operadoras 4G europeias aceitam algoritmos nulos em sessões não emergenciais, e que falhas de implementação em certos dispositivos permitem o mesmo comportamento do lado do UE.

Access Reject: A FBS transmite com potência superior às células legítimas, atraindo o UE e rejeitando as tentativas de conexão durante a fase de pré-autenticação. O atacante pode também configurar campos de barramento de acesso nas mensagens de *System Information*, forçando o UE a aguardar até 300 segundos antes de reavaliar a célula. Como o UE busca apenas a célula mais forte em cada frequência, a falha na conexão com essa célula impede-o de encontrar a estação base legítima na mesma banda.

Handover e Reseleção: A exploração de procedimentos de *handover* (UE conectado) ou reseleção de célula (UE em modo *idle/inactive*) permite forçar o UE a migrar para uma célula maliciosa. Para o *handover*, a FBS deve replicar parâmetros da célula legítima, incluindo identificador de célula, identidade da área de rastreamento, frequência de *downlink* e mensagens de informação de sistema. A reseleção explora o mesmo princípio, porém direciona-se a UEs em modo ocioso ou inativo. Ambos servem frequentemente como etapa preparatória para ataques subsequentes, como MiTM ou *International Mobile Subscriber Identity* (IMSI) *catching*.

IMSI Catcher: Força o UE a revelar identificadores permanentes ou temporários. É especialmente eficaz em 2G–4G, onde o IMSI pode ser transmitido em texto claro. Múltiplas mensagens de *downlink* em cada geração podem forçar a exposição do IMSI ou a exclusão do *Temporary Mobile Subscriber Identity* (TMSI). No 5G, embora o *Subscription Concealed Identifier* (SUCI) impeça a obtenção direta do identificador permanente, variantes conhecidas como *SUCI catchers* permitem ao adversário confirmar a presença de um UE específico, desde que já conheça um SUCI previamente utilizado ou o IMSI do alvo.

Ataques Passivos: O adversário apenas monitora a interface aérea, coletando mensagens *broadcast*, *paging* e sinalização inicial, sem realizar nenhuma transmissão. Essas informações são fundamentais para reconhecimento, *fingerprinting* de células/UEs e preparação de ataques ativos subsequentes. Este método é efetivamente indetectável e impõe demandas mínimas de energia e processamento.

8.3 Tipos de Ataque

Os tipos de ataque descrevem *o que* o adversário busca alcançar, independentemente do método utilizado.

Denial of Service (DoS): Visa impedir que o UE acesse ou mantenha conectividade com a rede. Em redes 5G, o atacante pode responder a solicitações de registro com rejeições contendo causas que impedem nova busca, ou enviar mensagens RRC malformadas que provocam falhas de software. Em redes 4G, técnicas incluem o envio de solicitações falsas para dessincronizar números de sequência entre UE e núcleo, a implantação de duas FBS que alternam rejeições e redirecionamentos aprisionando o UE em ciclo, e o esgotamento do limite de conexões RRC de uma estação base legítima.

Downgrade Tecnológico: Força o UE a migrar para tecnologias legadas com mecanismos de segurança mais fracos. Em 5G, uma abordagem comum envia rejeição de registro indicando que o modo 5G não é permitido; em 4G, a rejeição de atualização de área de rastreamento (*Tracking Area Update*) com causa de falha severa provoca efeito análogo. Avaliações empíricas demonstraram que todos os dispositivos comerciais testados eram vulneráveis a esses ataques. Outras técnicas incluem reprodução de mensagens de autenticação com números de sequência inválidos e envio de RRC *Release* com redirecionamento para tecnologia inferior.

Disrupção de Serviço: Introduz atrasos, perdas de pacotes ou degradação de *Quality of Service* (QoS), impactando aplicações sensíveis à latência. Um exemplo concreto explora a solicitação de funcionalidades (ou especificações) do UE antes do estabelecimento de segurança RRC: mediante MiTM, o adversário modifica a mensagem de funcionalidades, reduzindo a categoria do dispositivo, desabilitando múltiplas antenas, agregação de portadoras e voz sobre 4G, e restringindo bandas de frequência. Como as funcionalidades modificadas ficam armazenadas no núcleo da rede, o UE reconecta-se com configuração degradada.

Spoofing de Mensagens: Inclui o envio de SMS fraudulentos, comum após *downgrade* para 2G, e a falsificação de *Public Warning System* (PWS). A ausência de integridade em mensagens *broadcast* permite campanhas de desinformação e fraude. Atacantes também podem impedir que UEs em modo *idle* recebam alertas legítimos ao sequestrar o canal de *paging*.

Espionagem: Interceptação de comunicações possibilitada por *downgrades* criptográficos ou aceitação de algoritmos nulos. Em 4G, é possível contornar o procedimento de autenticação após reseleção de célula, estabelecendo o túnel de dados sem *ciphering*. Em 5G, técnicas de *fingerprinting* de tráfego cifrado permitem identificar *websites* ou vídeos acessados mesmo sem acesso ao conteúdo em texto claro.

Rastreamento: Explora identificadores permanentes ou temporários e padrões de *paging* para inferir presença e mobilidade do usuário. Em 5G, o atacante pode descartar mensagens de atualização de identificador temporário via MiTM, tornando o UE rastreável. Em 4G, correlação entre identificadores temporários pode ser obtida por monitoramento passivo. Em 2G, a ausência de proteção em relatórios de medição permite trilateração da posição do UE.

Tampering: Em redes onde a proteção de integridade do plano do usuário não está habilitada, atacantes podem modificar pacotes em trânsito. Um ataque demonstrado na prática

altera requisições *Domain Name System* (DNS) em pacotes cifrados, redirecionando o UE a um servidor malicioso sem quebrar o *ciphering*.

Sequestro de Sessão: O atacante assume o controle de uma sessão legítima, podendo acessar serviços vinculados ao endereço *Internet Protocol* (IP) do UE ou realizar atividades ilícitas atribuídas à vítima.

Fraude de Cobrança: Explora personificação e sequestro de sessão para consumir serviços faturados ao usuário legítimo. Uma variante utiliza MiTM distribuído para forçar conexão a uma célula em *roaming*, gerando cobranças excessivas.

Dreno de Energia: Força o UE a permanecer ativo por meio de *paging* falso, reconexões frequentes ou falhas repetidas de registro, sendo especialmente danoso para IoT. Em 4G, um ataque via MiTM contra dispositivos IoT remove o parâmetro de temporizador ativo durante o registro, desabilitando o modo de economia de energia e mantendo consumo elevado até que nova atualização de área (*Tracking Area Update*) ou registro seja realizado.

8.4 Detecção e Mitigação

8.4.1 Detecção

A detecção de FBS pode ser classificada em quatro modelos: baseada no UE, estacionária, baseada na rede e híbrida.

No UE, aplicações podem analisar potência de sinal (valores acima de -40 dBm são fortes indicativos de FBS), parâmetros inválidos de célula, localização implausível e velocidade anômala de *handover*. Uma implementação na China incluiu mais de 90 milhões de usuários, estimando a posição da FBS com erro mediano de 11 metros [72]. No iOS, aplicações baseadas em engenharia reversa da interface de banda base comparam características da célula contra bases de dados de células legítimas [73]. Abordagens nativas no *firmware* oferecem maior flexibilidade, permitindo análise direta da pilha de protocolos; propostas recentes utilizam a duração do estabelecimento de conexão e a contagem de tentativas repetidas como indicadores [74], reduzindo falsos positivos.

Dispositivos estacionários ou em veículos permitem monitoramento contínuo do ambiente de rádio. Múltiplos sensores possibilitam construir um mapa da infraestrutura celular e identificar estações anômalas [75]. Técnicas de *fingerprinting* de radiofrequência associam transmissões a dispositivos físicos específicos com precisão de até 99%, facilitando investigação forense [76]. Na rede, indicadores como mudanças inesperadas de área de localização, *downgrades* de algoritmos criptográficos e tempos anômalos de autenticação podem revelar a presença de FBS [77].

Abordagens híbridas combinam informações do UE e da rede, aumentando acurácia. Uma abordagem baseada em relatórios de medição dos UEs foi avaliada em rede comercial e incluída como anexo informativo na especificação 3GPP *Technical Specification* (TS) 33.501 [78]. Extensões utilizam valores de potência de sinal para estimar a posição da FBS por trilateração, alcançando precisão métrica em simulação [79]. Técnicas de aprendizado de máquina aplicadas a características de potência de sinal mostraram-se promissoras para classificação automática de células falsas [80].

8.4.2 Mitigação

A mitigação complementa a detecção, atuando para interromper um ataque em curso ou reduzir seu impacto. A desconexão imediata da FBS é a estratégia mais direta, mas não impede reconexão, devendo ser combinada com mecanismos adicionais.

O roteamento baseado em reputação mantém um registro de estações base previamente validadas, permitindo ao UE selecionar uma célula confiável após detecção de FBS. A seleção aleatória entre múltiplas estações na lista de reputação reduz a probabilidade de reconexão a uma FBS que falsifique identificadores legítimos [81].

Listas de bloqueio impedem conexões a uma estação detectada como falsa, embora a falsificação de identidade de células legítimas possa causar bloqueio incorreto de estações válidas [82]. No caso de *spoofing* de mensagens, a marcação como spam das mensagens de FBS protege o usuário final, movendo-as automaticamente para uma pasta de spam com alertas visuais.

8.5 Boas Práticas de Segurança

A análise realizada evidencia um desalinhamento recorrente entre recomendações do 3GPP e implantações reais. Avaliações empíricas revelaram operadoras 4G aceitando algoritmos nulos em sessões não emergenciais e dispositivos com falhas de implementação que permitem o mesmo comportamento. Em redes 2G, algoritmos criptográficos vulneráveis podem ainda estar em uso. As boas práticas recomendadas incluem:

1. **Algoritmos de integridade:** Rejeitar algoritmos nulos de integridade para sinalização de controle fora de cenários de emergência, tanto na rede quanto no UE.
2. **Ciphering:** Habilitar *ciphering* para tráfego de controle e de dados sempre que possível. Em redes 2G, utilizar algoritmos mais robustos sempre que suportados.
3. **Integridade do plano do usuário:** Em redes 5G, avaliar os casos de uso ao decidir sobre ativação. Muitas aplicações não seriam substancialmente impactadas por aumentos de latência de até 40%, tornando essa proteção viável.
4. **Configuração do 2G:** Evitar *fallback* automático para 2G quando existirem alternativas com autenticação mútua. Fabricantes devem considerar desabilitar o suporte a 2G.
5. **Ordem de consulta de funcionalidades:** Evitar solicitar funcionalidades do UE antes da ativação de segurança do RRC. Avaliações em 30 redes 4G de 20 países revelaram que 20 operadoras solicitavam funcionalidades antes da ativação de segurança.
6. **Proteção do SUCI:** Em redes 5G, garantir que esquemas de proteção não-nulos e chaves públicas estejam provisionados nos usuários.
7. **Transparência ao usuário:** Fornecer indicações sobre o estado de segurança da conexão. O sistema Android introduziu suporte a notificações quando o UE se conecta a rede sem *ciphering* ou transmite identificadores permanentes, embora desabilitadas por padrão.

8.6 Correções e Recomendações para o 6G

A raiz das vulnerabilidades reside na sinalização pré-autenticação desprotegida e na natureza opcional de mecanismos críticos. A padronização do 6G oferece a oportunidade de incorporar mecanismos de segurança como requisitos mandatórios desde a fase de projeto.

R1 — Autenticar *broadcast* e sinalização pré-autenticação. Exigir autenticação e integridade para todas as mensagens *broadcast* e pré-autenticação, incluindo blocos de informação de sistema, *paging* e sinalização RRC/NAS inicial. Abordagens existentes baseadas em

esquemas hierárquicos de identidade demonstram viabilidade, com latência criptográfica na faixa de microssegundos para assinatura e abaixo de um milissegundo para verificação.

R2 — Autenticar blocos de transporte. Incluir mecanismos que detectem ataques MiTM por meio de *Cyclic Redundancy Check* (CRC) criptográfico que vincule cada transmissão aos recursos de tempo-frequência utilizados, impedindo retransmissão transparente por *relay* malicioso. Não introduz sobrecarga de dados, mas acarreta custo computacional por bloco e não é compatível com repetidores transparentes legítimos.

R3 — Integridade do plano do usuário por padrão. A natureza opcional no 5G expõe usuários a adulteração de dados, sequestro de sessão e personificação. Avaliações indicam aumento de latência de 25% a 40%, motivando otimizações como aceleração por *hardware* e primitivas criptográficas leves.

R4 — Proibir *fallback* para tecnologias sem autenticação mútua. Todos os procedimentos de redirecionamento entre tecnologias devem ser criptograficamente protegidos; UEs devem aceitar comandos de *fallback* somente quando autenticados e com integridade verificada.

R5 — Detecção nativa de FBS na arquitetura. O núcleo da rede deve incorporar análise em tempo real de relatórios de medição, potencialmente com técnicas de inteligência artificial. Indicadores de segurança no UE devem ser padronizados e habilitados por padrão.

R6 — Proteger informações sensíveis antes do contexto seguro. Funcionalidades do UE, identificadores de equipamento e dados de medição devem ser transmitidos somente após autenticação e acordo de chaves.

Essas recomendações devem ser implementadas por meio de perfis de segurança adaptativos: proteções para procedimentos críticos de controle constituem linha de base mandatória, enquanto mecanismos adicionais podem ser aplicados conforme restrições de latência, energia e cobertura. Mecanismos de integridade devem preservar compatibilidade com implantações baseadas em *relay*, e opções de autenticação leve devem ser suportadas para dispositivos IoT.

8.7 Conclusão

Ataques baseados em FBS continuam sendo uma ameaça prática e relevante. Nenhuma abordagem isolada de detecção é suficiente: métodos baseados no UE oferecem proteção rápida porém com visibilidade limitada; soluções de rede e híbridas aproveitam maior consciência contextual ao custo de maior complexidade; e dispositivos estacionários complementam com monitoramento contínuo e suporte forense.

A mitigação efetiva requer uma abordagem holística combinando detecção, configurações seguras e integração nativa de mecanismos de segurança. À medida que as redes evoluem para o 6G, esses mecanismos devem ser mandatórios, garantindo proteção robusta sem comprometer desempenho e flexibilidade.

9 Estratégias de Uplink NOMA Baseadas em FreeChirp para Redes IoT de Comunicação Direta ao Satélite

*Felipe Augusto Tondo, Oana Iova, Miguel Gutiérrez Gaitán, Samuel Montejo Sánchez,
Richard Demo Souza*
UFSC

9.1 Introdução

A contínua evolução das tecnologias sem fio tem impulsionado o crescimento do número de dispositivos de IoT conectados, intensificando também a demanda por cobertura global [83]. No entanto, as redes 5G ainda não atendem plenamente aos requisitos de conectividade massiva em escala global [84], não havendo garantias de que tais demandas sejam supridas mesmo com sua expansão. Nesse contexto, os sistemas 6G buscam superar limitações associadas à vida útil dos dispositivos, aos custos de implementação, à confiabilidade e à complexidade de *hardware* [85]. Assim, a comunicação espacial baseada em CubeSats destaca-se como uma alternativa promissora para viabilizar a conectividade ubíqua no 6G [86].

Para o cenário onde os dispositivos transmitem informações diretamente ao satélite, conhecido como DtS-IoT, os satélites do tipo *Low Earth Orbit* (LEO) destacam-se devido ao menor custo, à modularidade e à menor latência quando comparados aos satélites em órbita média e geoestacionária [87]. Além disso, esses satélites apresentam menores perdas de propagação e potencial para prover cobertura global por meio de constelações de satélites, consolidando-se como elementos relevantes para sistemas 5G e 6G [88]. Atualmente, companhias como a Lacuna Space e a Myriota utilizam a tecnologia LoRaWAN como base para o fornecimento de serviços IoT em aplicações como agricultura inteligente, mineração e monitoramento de equipamentos pesados [89, 90].

Apesar dos avanços, o DtS-IoT ainda apresenta desafios consideráveis devido ao tempo limitado de cobertura, à forte dependência do ângulo de elevação e às janelas de visibilidade semelhantes [91]. Uma das principais dificuldades está no desenvolvimento de protocolos eficientes para cenários de alta densidade. Esquemas simples, como o ALOHA puro, são amplamente utilizados, mas sofrem com congestionamentos e colisões à medida que a rede cresce. Alternativas baseadas em *slotted* ALOHA ou TDMA [92, 93, 94, 95] reduzem colisões, porém exigem sincronização rigorosa ou conhecimento completo da rede, limitando sua aplicação prática. Para contornar esse desafio em redes DtS-IoT baseadas em LoRaWAN, um estudo recente propõe um protocolo denominado *Free Signal Multiple Access* (FSMA), que utiliza um sinal de controle chamado FreeChirp, permitindo que os dispositivos realizem o processo de *Channel Activity Detection* (CAD) e transmitam apenas quando o canal estiver livre. Contudo, em cenários densos, o FSMA ainda apresenta elevada contenção de canal, comprometendo sua escalabilidade.

Com o objetivo de mitigar o problema de colisões e aumentar o uso eficiente da cobertura em DtS-IoT, os autores em [96] propuseram novos algoritmos que permitem que dispositivos IoT estimem as janelas de visibilidade de satélites. Com base nesse contexto, os autores em [97] introduziram estratégias avançadas para DtS-IoT usando NOMA, com FTP ou CTP. Apesar dos ganhos em termos de quantidade de informação exitosa recebida no satélite, as estratégias apresentam um decréscimo à medida que muitos dispositivos possuem oportunidades de transmissão semelhantes.

Diferentemente de trabalhos anteriores [98, 97], este capítulo propõe duas novas estratégias de transmissão para o cenário DtS-IoT, integrando o sinal FreeChirp — como mecanismo de

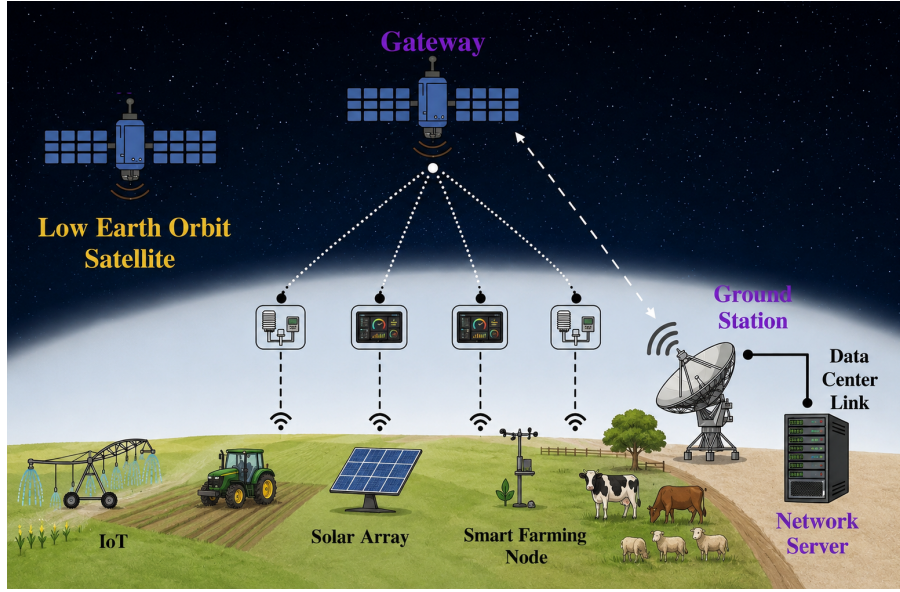


Figura 21: Perspectiva geral de uma rede de comunicação IoT direta ao satélite.

controle por parte do satélite — com as estratégias NOMA-FTP e NOMA-CTP, com o objetivo de reduzir colisões e aumentar a eficiência da janela de visibilidade.

9.2 Metodologia

Esta seção é responsável por apresentar a arquitetura do sistema, o modelo de canal para o cenário DtS-IoT, e as novas estratégias de transmissão.

Arquitetura do Sistema

Para o modelo do sistema, a Figura 21 apresenta um cenário em que os dispositivos estão distribuídos ao longo de uma área de interesse e utilizam a tecnologia LoRa na modulação do sinal para transmitir dados a um satélite LEO, equipado com um gateway com suporte ao cancelamento sucessivo de interferência. Em aplicações de monitoramento IoT, os dados recebidos pelo satélite são encaminhados para uma estação terrestre e, posteriormente, para um servidor central.

Modelo de Canal para Enlace DtS-IoT

O modelo de canal entre o dispositivo e o satélite assume uma componente determinística, relacionada à perda de percurso e à órbita, e uma componente aleatória, associada ao desvanecimento do enlace. Considerando a transmissão de um dispositivo de referência $x \in \{1, 2, \dots, \mathcal{X}\}$, o sinal recebido no satélite corresponde à soma do sinal útil atenuado, da interferência gerada pelo conjunto I contendo os demais dispositivos que transmitem simultaneamente com o dispositivo x , e do ruído do canal, conforme descrito a seguir:

$$r = \sqrt{P_x g_x} h_x s_x + \sum_{i \in I} \sqrt{P_i g_i} h_i s_i + w, \quad (53)$$

em que P representa a potência de transmissão, g o ganho do canal associado aos efeitos de larga escala, h os efeitos de *fading* e s o sinal modulado. Os parâmetros indexados por i

possuem o mesmo significado para os dispositivos interferentes. O termo w representa o ruído branco gaussiano aditivo (AWGN), com média zero e potência $\sigma_w^2 = -174 + F + 10 \log_{10} B$ dBm, considerando fator de ruído F dB e largura de banda B Hz [99].

A comunicação ocorre apenas durante a janela de visibilidade, quando o ângulo de elevação é maior que $\alpha_{\min}$. A distância entre o dispositivo IoT e o satélite é calculada por $d(\alpha) = [\sqrt{(R+H)^2 - (R \cos(\alpha))^2} - R \sin(\alpha)]$, onde α é dado em graus, $R = 6,378 \times 10^6$ m é o raio da Terra e H é a altitude do satélite. O ganho de canal em larga escala entre o dispositivo x e o satélite é modelado usando a equação de propagação em espaço livre dada por $g = G_t G_r (\lambda/4\pi d)^2$ [100], na qual λ corresponde ao comprimento de onda, e G_t e G_r representam os ganhos das antenas transmissora e receptora, respectivamente.

Para modelar o desvanecimento sofrido pelo enlace DtS-IoT, assumiu-se a combinação de dois processos aleatórios [97, 101, 102], conforme expressado por

$$f_h(h) = \underbrace{\int_0^\infty 2(K+1) \frac{h}{S^2} \frac{e^{-K}}{e^{(K+1)\frac{h^2}{S^2}}} I_0\left(2\frac{h}{S} \sqrt{K(K+1)}\right) dS}_{f_h(h|S)} \times \underbrace{\frac{1}{\sqrt{2\pi}\beta\sigma} e^{-\frac{1}{2}\left(\frac{\ln S - \mu}{\beta\sigma}\right)^2}}_{f_S(S)} dS, \quad (54)$$

em que $f_h(h|S)$ representa a função densidade de probabilidade Riceana condicionada ao *shadowing* S . O fator de Rice K varia de acordo com α , de modo que maiores elevações resultam em condições de visada mais favoráveis [101]. Os parâmetros μ e $(\beta\sigma)^2$ representam a média e a variância da componente log-normal associada ao *shadowing*. Nas simulações, os parâmetros do canal são definidos com base em valores empíricos dependentes do ângulo de elevação reportados em [101].

Estratégias NOMA baseadas em FreeChirp

De forma simplificada, o FreeChirp consiste em um LoRa *up-chirp* transmitido periodicamente pelo satélite com um intervalo de $t_{\text{fchirp}} = t_{\text{chirp}} + t_{\text{wait}}$, onde t_{chirp} denota a duração do sinal FreeChirp. O satélite transmite o sinal de controle sempre que o canal está livre. Em seguida, aguarda um período predefinido t_{wait} para verificar uma possível transmissão. Ao detectar uma transmissão, o gateway suspende a emissão de novos FreeChirps até o término da recepção.

Do lado dos dispositivos, assume-se o conhecimento prévio das janelas de visibilidade do satélite por meio de mecanismos de estimação orbital da trajetória do satélite [103, 97, 104]. O dispositivo IoT acorda durante a passada do satélite e realiza o CAD durante um intervalo t_{sense} . Com base na escuta do canal e na detecção do FreeChirp, o dispositivo procede assim: (i) se um FreeChirp for detectado, o dispositivo inicia a transmissão imediatamente; (ii) caso contrário, o dispositivo entra em modo de repouso (*sleep*) e não utiliza essa oportunidade de transmissão.

Inspiradas nesse contexto, são propostas duas novas estratégias de transmissão para DtS-IoT, nas quais o mecanismo de controle baseado em FreeChirp é combinado com estratégias NOMA de potência de transmissão fixa (FChirp-FTP) e potência de transmissão controlada (FChirp-CTP). Assim como na estratégia original [97], no FChirp-FTP os dispositivos selecionam instantes específicos dentro da janela de visibilidade capazes de gerar níveis projetados de potência recebida no satélite, denotados por $\mathcal{L}_1$ e $\mathcal{L}_2$. Já no FChirp-CTP, os dispositivos podem transmitir em qualquer instante, ajustando sua potência de transmissão conforme necessário. Em ambos os casos, o FreeChirp atua como mecanismo de coordenação para reduzir

colisões e permitir transmissões simultâneas por meio do processo de cancelamento sucessivo de interferência no receptor.

Além disso, os esquemas FSMA [98] e FChirp-ALOHA são usados como bases de comparação, aplicando-se o FreeChirp, porém sem NOMA. Adicionalmente, assumiu-se o efeito de captura do LoRa, permitindo a decodificação do sinal mais forte quando a diferença de potência recebida ultrapassa um limiar predefinido [105].

Condições de Decodificação e Projeto dos Níveis NOMA

Seguindo alguns estudos sobre tecnologia LoRaWAN aplicado em DtS-IoT [106, 107, 108], assumimos que duas condições devem ser atendidas para garantir a decodificação bem-sucedida no receptor [99]:

$$\begin{cases} \mathcal{C}_1 : \text{SNR}_x = \frac{P_x g_x h_x^2}{\sigma_w^2} \geq \gamma \\ \mathcal{C}_2 : \text{SIR}_x = \frac{P_x g_x h_x^2}{\sum_{i \in I} P_i g_i h_i^2} \geq \psi. \end{cases} \quad (55)$$

Dessa forma, a primeira condição estabelece que a relação sinal-ruído (SNR) no gateway deve estar acima de um limiar γ definido pela tecnologia LoRa [99]. Portanto, a potência recebida deve ser suficientemente elevada para que o gateway detecte e decodifique a mensagem acima do nível de ruído. A segunda condição estabelece que a relação sinal-interferência (SIR) no receptor deve estar acima do limiar de captura ψ , o qual é fixo para uma determinada tecnologia e amplamente estudado no contexto LoRaWAN [105]. Nesse caso, a potência recebida de uma transmissão deve ser suficientemente maior que a soma das interferências presentes no canal.

Para o projeto de $\mathcal{L}_1$, assume-se que um dispositivo não consegue estabelecer comunicação quando a SNR recebida no satélite é inferior ao limiar γ [109], de modo que a probabilidade de *outage* é dada por $P_{\text{out, SNR}} = \Pr(\mathcal{L}_1 h_x^2 / \sigma_w^2 < \gamma)$. Assim, $\mathcal{L}_2$ é definido considerando a presença de interferência e um limiar mínimo ψ necessário para garantir decodificação via cancelamento sucessivo, resultando em $P_{\text{out, SIR}} = \Pr(\mathcal{L}_2 h_x^2 / \mathcal{L}_1 h_1^2 < \psi)$. Considerando $\alpha_{\min} = 30^\circ$, $\sigma_w^2 = -117$ dBm e uma restrição de confiabilidade de $P_{\text{out, SNR}} < 10\%$, obteve-se $\mathcal{L}_1 = -124.40$ dBm. De forma semelhante, o segundo nível foi definido como $\mathcal{L}_2 = -121.50$ dBm, ambos obtidos a partir de extensas simulações numéricas em [2].

Parâmetros Temporais FreeChirp

Após cada transmissão do sinal FreeChirp, é necessário um intervalo de espera para permitir a detecção adequada de pacotes no gateway. Conforme especificado pela documentação do transceptor Semtech SX126x [110], a detecção confiável do preâmbulo requer a observação de 4 a 8 símbolos. Assim, seguindo uma abordagem conservadora semelhante à utilizada em [98], define-se

$$t_{\text{wait}} = 6 \cdot (t_{\text{sym}}) = \frac{6 \cdot (2^{\text{SF}})}{BW}, \quad (56)$$

em que t_{sym} representa a duração do símbolo LoRa para o fator de espalhamento (SF) e a largura de banda adotados. Diferentemente de um pacote LoRa convencional, o FreeChirp consiste em um sinal de excitação curto, de duração t_{chirp} , suficiente para acionar o CAD nos dispositivos.

Tabela 3: Parâmetros de Simulação do Cenário DtS-IoT com FreeChirp.

Parâmetros dos Dispositivos em Solo			Parâmetros do Satélite/Gateway		
Descrição	Valor	Parâmetro	Descrição	Valor	Parâmetro
Potência Máxima de Transmissão	14 dBm	$P_{\max}$	Ganho da Antena Receptora	13.5 dBi	G_r
Ganho da Antena Transmissora	0 dBi	G_t	Sensibilidade do Receptor	-137 dBm	-
Menor Nível de Potência Recebida	-124.40 dBm	$\mathcal{L}_1$	Limite de SNR	-20 dB	γ
Maior Nível de Potência Recebida	-121.50 dBm	$\mathcal{L}_2$	Limite de SIR	1 dB	ψ
Tamanho do Pacote	20 Bytes	b	Figura de Ruído	6 dB	F
Fator de Espalhamento	12	SF	Largura de Banda	125 kHz	B
Tempo no Ar	1.8104 s	ToA	Frequência da Portadora	868 MHz	f
Parâmetros FreeChirp					
Duração do FreeChirp	$t_{\text{sym}}/2$ s	t_{chirp}	Tempo de Espera do Gateway	$6 \cdot t_{\text{sym}}$ s	t_{wait}
Intervalo do FreeChirp	$t_{\text{chirp}} + t_{\text{wait}}$ s	t_{fchirp}	Período de Escuta	$\tau \cdot t_{\text{chirp}}$ s	t_{sense}

9.3 Resultados

A seguir, são apresentados resultados numéricos para os esquemas propostos FChirp-FTP e FChirp-CTP, considerando os parâmetros da Tabela 3. Os esquemas são comparados com FSMA [98], FChirp-ALOHA e NOMA FTP/CTP [2], permitindo avaliar os ganhos da coordenação FreeChirp e do domínio de potência NOMA. O desempenho é analisado por meio do *goodput*, definido como o número médio de bytes recebidos com sucesso no satélite por volta (referida nesse capítulo como *lap*), dado por $\mathcal{G} = \mathcal{M}b$, em que $\mathcal{M} \leq U$ representa o número de mensagens recebidas com sucesso e b o número de bytes por mensagem.

Os dispositivos IoT estão distribuídos uniformemente ao longo da França, usando coordenadas geográficas reais geradas pela biblioteca GeoPy [111]. A biblioteca Skyfield [112] é usada para calcular a distância entre os nós terrestres e o gateway, enquanto os dados orbitais são coletados do CelesTrack no formato de dois elementos de linha [113]. Para apresentar um cenário mais prático, as simulações⁸ consideram a trajetória real do LacunaSat-3, orbitando a altitudes de 500 a 600 km acima da superfície [114]. Seguindo [103, 97], adota-se o fator de espalhamento mais robusto SF = 12. O desempenho é avaliado para um número total de dispositivos $\mathcal{X} \in \{2, 25, 50, \dots, 2000\}$, com um passo de cinquenta dispositivos para $\mathcal{X} \geq 50$.

Taxa de Transferência Efetiva

A Figura 22 apresenta o número médio de bytes recebidos com sucesso por passada para as estratégias FChirp-FTP e FChirp-CTP, comparadas aos esquemas NOMA originais [97], FChirp-ALOHA e FSMA [98]. Os resultados demonstram ganhos expressivos de desempenho, especialmente em cenários densos. Os maiores valores de *goodput* foram obtidos por FChirp-FTP e FChirp-CTP, atingindo 2637.60 e 2452.70 bytes/lap para 900 e 1150 dispositivos, respectivamente, representando ganhos de até 169.41% e 200.57% em relação ao FSMA. Em cenários de baixa densidade ($\mathcal{X} \leq 200$), o FSMA apresenta melhor desempenho devido ao mecanismo de *backoff*, que reduz a probabilidade de colisões. Para um cenário denso com $\mathcal{X} = 2000$, o FChirp-FTP atinge 2306.60 bytes/lap (315.08% de melhoria) e o FChirp-CTP atinge 2200.60 bytes/lap (296.01% de melhoria) em relação aos esquemas originais NOMA FTP e CTP. No mesmo cenário, os ganhos em relação ao FSMA (571.2 bytes/lap) atingem 303.81% e 285.26%, respectivamente. Esses ganhos decorrem da combinação entre o mecanismo FreeChirp contro-

⁸O simulador foi adaptado para considerar os cenários deste trabalho, e a versão correspondente estará disponível publicamente em breve no GitHub (<https://github.com/tondoFe/LEONIS.git>).

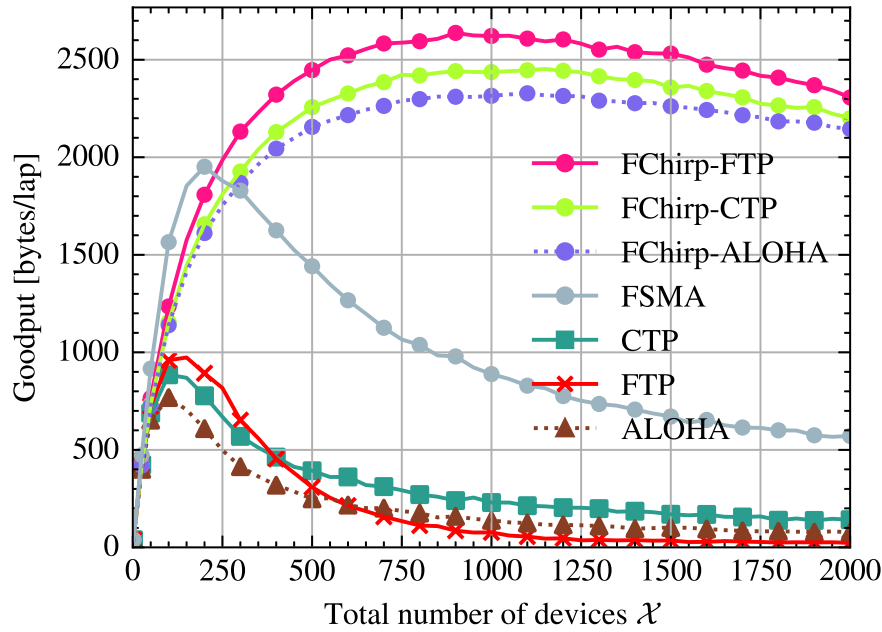


Figura 22: Número médio de bytes recebidos com sucesso por volta *versus* $\mathcal{X}$ para FChirp-FTP, FChirp-CTP, FChirp-ALOHA e FSMA, em comparação com os protocolos originais [2].

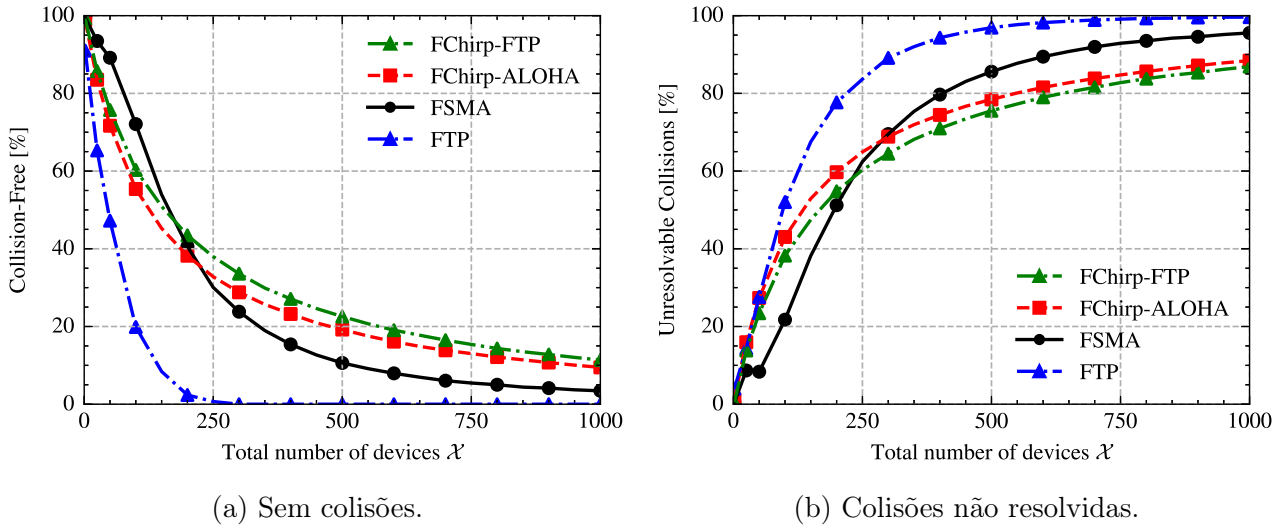


Figura 23: Percentagem de dispositivos associados a transmissões (a) sem colisões e (b) com colisões não resolvidas *versus* $\mathcal{X}$ para as estratégias FChirp-FTP, FChirp-ALOHA, FSMA e FTP.

lado pelo gateway e o uso de níveis de potência planejados, permitindo melhor aproveitamento temporal da janela de visibilidade do satélite e maior eficiência do processo de cancelamento sucessivo no gateway pelo planejamento adequado dos níveis de potência recebidos. Além disso, diferentemente do FSMA, os esquemas propostos evitam múltiplas tentativas de sensoriamento durante períodos de canal ocupado, reduzindo colisões e melhorando o aproveitamento das oportunidades de transmissão.

Análise do Cenário de Colisões

Para compreender melhor os resultados de *goodput*, apresenta-se na Figura 23a o percentual de dispositivos que realizam transmissões livres de colisão em função de $\mathcal{X}$. Em cenários esparsos ($\mathcal{X} < 150$), o FSMA obtém o melhor desempenho devido ao mecanismo de *backoff*, que redistribui as transmissões no tempo e reduz a probabilidade de colisões. Entretanto, à medida que o número de dispositivos aumenta, a estratégia FChirp-FTP passa a manter o maior percentual de transmissões sem colisão. Esse ganho é resultado da combinação entre o sensoriamento baseado em FreeChirp, que evita transmissões durante recepções em andamento, e a distribuição espacial das transmissões proporcionada pelo FTP. Em contraste, FSMA e FChirp-ALOHA apresentam redução mais acentuada no número de transmissões livres de colisão, pois dependem apenas da separação temporal e do efeito de captura, sem explorar a multiplexação no domínio de potência oferecida pelo NOMA. Já o FTP convencional apresenta a degradação mais rápida, uma vez que dispositivos que selecionam os mesmos instantes fixos de transmissão inevitavelmente colidem devido à ausência de mecanismos de coordenação.

A Figura 23b corrobora a ideia de que a estratégia FChirp-FTP apresenta o menor percentual de colisões que não conseguiram ser resolvidas. Enquanto FSMA e FChirp-ALOHA sofrem um crescimento acentuado dessas colisões com o aumento da densidade da rede, o FChirp-FTP consegue mitigar esse comportamento ao combinar o uso do sinal de controle e NOMA. Esse percentual de colisões não resolvidas é o principal fator responsável pelos elevados ganhos de *goodput* observados na Figura 22, especialmente em cenários densos, nos quais o FChirp-FTP alcança ganhos superiores a 315% em relação ao FTP convencional e acima de 300% quando comparado ao FSMA.

9.4 Conclusão

O capítulo apresentou novas estratégias de transmissão para o cenário DtS-IoT, combinando o mecanismo de acesso FreeChirp com técnicas NOMA no domínio da potência. As abordagens propostas, FChirp-FTP e FChirp-CTP, permitem melhor aproveitamento das janelas de visibilidade do satélite, buscando características essenciais como baixa complexidade, ausência de sincronização ou conhecimento completo do sistema, e comportamento agnóstico em relação à rede. Os resultados, obtidos com dados reais do satélite LacunaSat-3 sobre a França, demonstraram ganhos significativos de desempenho em termos de *goodput*, especialmente em cenários de alta densidade da rede.

Ressalta-se que novos estudos estão sendo conduzidos com o objetivo de projetar, de forma mais robusta, os parâmetros temporais utilizados pelo FreeChirp, buscando garantir uma melhor detecção. Além disso, o desempenho das estratégias também está sendo avaliado em relação à duração da escuta do canal pelos dispositivos no solo. Como limitação, destaca-se que o presente capítulo considera apenas um único satélite LEO e não incorpora análises detalhadas de eficiência energética associadas ao sensoriamento e à escuta do FreeChirp. Como trabalhos futuros, pretende-se investigar modelos de consumo energético realistas, cenários com múltiplos satélites e estratégias híbridas adaptativas que combinem características do FSMA e das abordagens NOMA propostas.

10 Otimização Conjunta Escalável de Recursos para Sistemas Cell-Free Massive MIMO Multi-CPU

Jonathan N. Gois
UFRJ

10.1 Introdução

Entre as tecnologias que viabilizam altas taxas de dados, arquiteturas com múltiplas antenas desempenham um papel de destaque. Um exemplo fundamental é o MIMO, que oferece benefícios como maior direcionalidade de feixe, aumento da taxa de transmissão e redução do consumo de potência, melhorando assim o desempenho global da rede [115]. Essas vantagens são ainda mais acentuadas quando o número de antenas transmissoras excede significativamente o número de UEs, como no *Massive Multiple-Input Multiple-Output* (mMIMO). No entanto, devido à persistente escassez de espectro, o mMIMO isoladamente não é suficiente para atender às crescentes demandas de taxa de dados das redes sem fio emergentes. Para contornar essa limitação, uma solução promissora é o *Cell-Free* (CF)-mMIMO [116], que distribui o arranjo de antenas entre diversos *Access Point* (AP)s e elimina o conceito tradicional de células. Ao invés do serviço a um dado terminal acontecer por um único ponto de acesso com um grande número de antenas, o CF-mMIMO propõe serviço fornecido por diversos pontos de acesso simultaneamente, explorando a diversidade espacial e cooperação entre os pontos de acesso.

A implementação puramente centralizada do CF-mMIMO requer uma *Central Processing Unit* (CPU), que atua como o núcleo de coordenação da rede distribuída de APs, realizando tarefas como coordenação entre APs, agregação de *Channel State Information* (CSI), alocação de recursos e sincronização do sistema. Embora uma CPU totalmente centralizada seja capaz de otimizar recursos globalmente, ela apresenta limitações em termos de escalabilidade, latência e robustez. Por outro lado, a distribuição do processamento entre múltiplas CPUs reduz a sobrecarga do *backhaul*, mas exige estratégias eficazes de cooperação para manter a consistência no processamento de dados, na alocação de potência e na sincronização.

Diversos autores propuseram modelos hierárquicos para caracterizar os níveis de cooperação entre CPUs em sistemas CF-mMIMO. Por exemplo, Kim *et al.* [117] definiram uma hierarquia de cooperação com quatro níveis. No nível mais baixo, **Nível 1**, os APs operam de forma totalmente distribuída, com cada AP estimando independentemente o CSI e determinando sua própria política de alocação de potência. No **Nível 2**, a CPU coordena grupos disjuntos de APs, sendo a maior parte do processamento realizada localmente dentro de cada grupo. O **Nível 3** introduz um grau de escalabilidade ao permitir que CPUs compartilhem informações limitadas sobre o ambiente, possibilitando uma coordenação parcial. Por fim, no **Nível 4**, ocorre uma troca abrangente de informações entre todas as CPUs. Esse nível mais alto de cooperação confere ao sistema uma visão global do ambiente, mas exige elevada capacidade de *backhaul* para suportar a carga de comunicação resultante.

Li *et al.* [118] propuseram uma estrutura hierárquica semelhante, utilizando o nível de troca de informações entre CPUs como principal critério de classificação. Nesse modelo, o **Nível 1** representa um cenário sem comunicação entre CPUs. Os **Níveis 2 e 3** envolvem a troca apenas de CSI estatístico, mas diferem na forma como essa informação é processada e distribuída, refletindo variações no grau de centralização e nas responsabilidades computacionais. No **Nível 4**, o CSI instantâneo é compartilhado entre todas as CPUs, proporcionando o maior nível de cooperação, porém ao custo de elevada demanda de *backhaul* e da necessidade de mecanismos

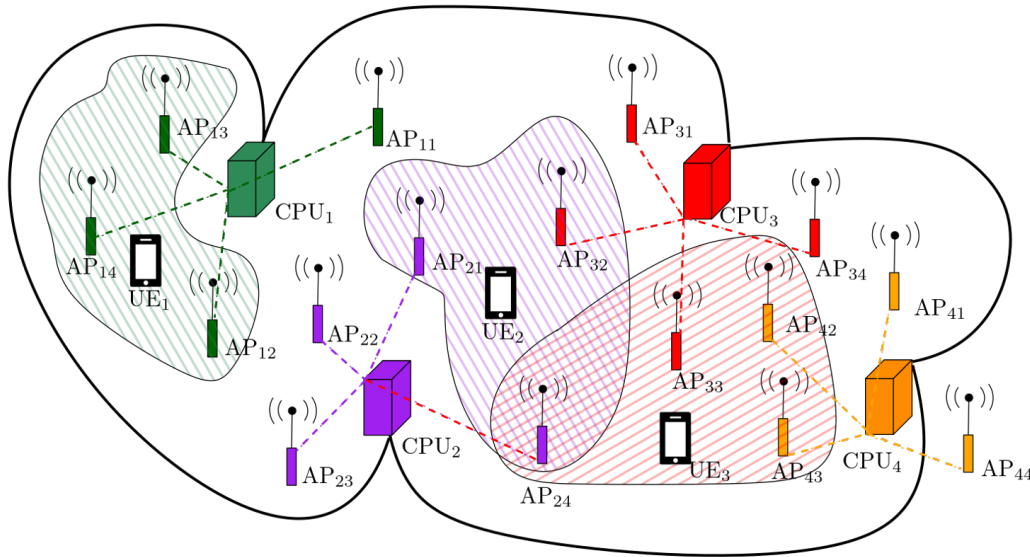


Figura 24: Arquitetura CF-mMIMO com múltiplas CPUs em sua versão escalável.

altamente eficientes de troca de dados.

A escalabilidade é sempre uma preocupação central no CF-mMIMO, tornando essencial compreender os desafios de colaboração para viabilizar implementações reais dessa tecnologia. Um avanço nessa direção é o *Dynamic Cooperation Clustering* (DCC) [119]. Idealmente, cada UE estaria conectado a todos os APs, mas esse cenário é impraticável. No entanto, o DCC explora o fato de que nem todos os APs contribuem igualmente para a *Spectral Efficiency* (SE) de um determinado enlace. Assim, a ideia é limitar o número de conexões por UE. No CF-mMIMO centrado no usuário, cada UE é atendido por um conjunto dedicado de APs, o que proporciona a melhor taxa de dados possível dentro de um conjunto restrito de APs. Os autores denominaram esta configuração como *Scalable Cell-Free Massive MIMO* e ilustraram tal sistema na Figura 24.

Essa solução, entretanto, traz consigo um custo computacional próprio. O CSI precisa ser atualizado com frequência, e o grupo de APs deve ser continuamente reconfigurado, tornando-se cada vez mais oneroso à medida que o número de APs e UEs cresce. A carga computacional associada à formação de *clusters* e à alocação de potência aumenta significativamente, agravando ainda mais os problemas de escalabilidade. O desafio se intensifica em sistemas multi-CPU sem cooperação, nos quais estratégias ingênuas, como a simples ordenação dos valores de CSI, não garantem uma clusterização eficiente. Para superar essas limitações, uma das propostas do presente trabalho é um algoritmo não supervisionado de baixa complexidade, projetado para aumentar a SE sem requerer cooperação entre CPUs. O método concentra o DCC dentro do domínio de uma mesma CPU, reduzindo assim a carga no *backhaul* e a complexidade computacional, ao mesmo tempo em que mantém o desempenho em termos de SE.

Ao passo que essa solução torna o sistema escalável no nível de formação de *clusters*, ainda permanecem gargalos importantes para uma implementação distribuída com carga sustentável no *backhaul*. Um deles é a alocação de potência em *downlink*. Sob o critério de justiça *max-min*, que busca maximizar a menor taxa entre os UEs, é necessário algum nível de coordenação entre CPUs para estimar corretamente o desempenho global dos usuários e redistribuir potência de forma justa. Assim, além da formação escalável de *clusters*, este trabalho também considera uma estratégia de refinamento de potência de baixa dimensão, voltada à melhoria da SE mínima com reduzido *overhead* de comunicação.

10.2 Modelo do Sistema

Considera-se um sistema CF-mMIMO em *downlink* composto por L APs, cada um equipado com N antenas, servindo K UEs com uma antena cada. Os APs são distribuídos aleatoriamente e de forma uniforme sobre uma área quadrada, adotando-se condição de contorno periódica (*wrap-around*) para reduzir efeitos de borda. Os UEs também são posicionados de forma aleatória e independente. Um exemplo de disposição dos elementos da rede é apresentado na Figura 25.

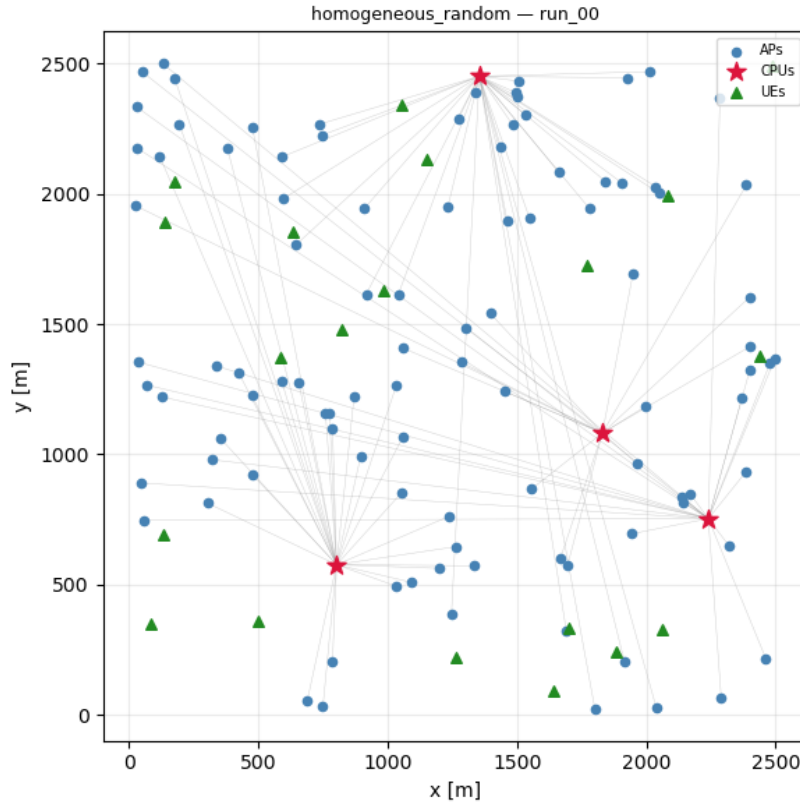


Figura 25: Exemplo de cenário utilizado no trabalho. As linhas indicam a associação entre cada CPU e os APs sob sua responsabilidade. O critério de associação é baseado na menor distância considerando *wrap-around*.

10.2.1 Topologia da Rede e Arquitetura Multi-CPU

O processamento do sistema é particionado em U CPUs. Cada CPU é responsável por um subconjunto disjundo de APs, de modo que [116]

$$\mathcal{L}_u \subseteq \{1, \dots, L\}, \quad \mathcal{L}_u \cap \mathcal{L}_{u'} = \emptyset \quad (u \neq u') \text{ e } \bigcup_{u=1}^U \mathcal{L}_u = \{1, \dots, L\}. \quad (57)$$

Dessa forma, cada AP é gerenciado por uma única CPU, enquanto um mesmo UE pode ser servido por APs pertencentes a diferentes CPUs. A associação entre APs, UEs e CPUs é

descrita pelo tensor binário

$$D[\ell, k, u] \in \{0, 1\}, \quad \sum_{u=1}^U D[\ell, k, u] \leq 1 \quad \forall \ell, k. \quad (58)$$

O elemento $D[\ell, k, u] = 1$ indica que o AP ℓ , gerenciado pela CPU u , serve ativamente o UE k . A restrição em (58) decorre do fato de que cada AP pertence a uma única CPU. A máscara global de associação AP-UE, independente da identificação da CPU, é definida como

$$D_{\text{flat}}[\ell, k] = \mathbf{1} \left[\sum_{u=1}^U D[\ell, k, u] > 0 \right] \in \{0, 1\}. \quad (59)$$

Essa máscara define quais APs transmitem dados para cada UE no *downlink*.

10.2.2 Modelo de Canal

O canal entre o AP ℓ e o UE k é modelado como um vetor aleatório complexo Gaussiano circularmente simétrico,

$$\mathbf{h}_{\ell k} \sim \mathcal{CN}(\mathbf{0}, \mathbf{R}_{\ell k}), \quad (60)$$

em que $\mathbf{R}_{\ell k} \in \mathbb{C}^{N \times N}$ é a matriz de correlação espacial. Essa matriz é escrita como

$$\mathbf{R}_{\ell k} = \beta_{\ell k} \mathbf{R}_{\text{local}}(\theta_{\ell k}, \sigma_{\text{ASD}}), \quad (61)$$

em que $\beta_{\ell k}$ representa o desvanecimento *Large Scale Fading* (LSF), $\theta_{\ell k}$ é o ângulo médio de chegada e σ_{ASD} é o espalhamento angular. O termo $\mathbf{R}_{\text{local}}(\cdot)$ modela a correlação espacial local entre as antenas do AP.

O coeficiente $\beta_{\ell k}$ agrega perda de trajeto (*path loss*) e *shadowing* lognormal:

$$\beta_{\ell k} [\text{dB}] = \underbrace{\text{PL}(d_{\ell k})}_{\text{perda de trajeto}} + \underbrace{f_{\ell k}}_{\text{shadowing}}, \quad (62)$$

em que $d_{\ell k}$ é a distância AP-UE (em metros) e $f_{\ell k}$ é o componente de *shadowing*. O presente trabalho usa (principalmente) o modelo de perda de percurso com três inclinações:

$$\text{PL}(d_{\ell k}) [\text{dB}] = \begin{cases} -L_1 - 35 \log_{10}(d_{\ell k}) & d_{\ell k} > d_1, \\ -L_1 - 15 \log_{10}(d_1) - 20 \log_{10}(d_{\ell k}) & d_0 < d_{\ell k} \leq d_1, \\ -L_1 - 15 \log_{10}(d_1) - 20 \log_{10}(d_0) & d_{\ell k} \leq d_0, \end{cases} \quad (63)$$

com limiares $d_0=10$ m e $d_1=50$ m, e constante de perda L_1 .

10.2.3 Dynamic Cooperation Clustering (DCC)

A arquitetura proposta com múltiplas CPUs introduz dois tipos de *clusters* com papéis distintos no sistema de comunicação implementado aqui:

1. Cluster Real

O *cluster* real da CPU u é o conjunto $\mathcal{L}_u$ de APs fisicamente conectados à CPU u (Eq. (57)). Os *clusters* reais formam uma partição disjunta e completa dos L APs.

2. Cluster Virtual

O *cluster* virtual do UE k é o conjunto de APs que transmite ativamente para o UE k no *downlink*:

$$\mathcal{V}_k = \{\ell : D_{\text{flat}}[\ell, k] = 1\}, \quad |\mathcal{V}_k| \leq L_{\text{max}}. \quad (64)$$

Os *clusters* virtuais não são disjuntos: um AP pode pertencer ao *cluster* virtual de múltiplos UEs simultaneamente.

Além disso, pode-se definir o conjunto de CPUs envolvidas na alocação de recursos de cada UE. Em outras palavras, para o UE k , define-se o conjunto de CPUs que possuem pelo menos um AP em seu *cluster* virtual:

$$\mathcal{Q}_k = \{u : \exists \ell \in \mathcal{L}_u \cap \mathcal{V}_k\} = \{u : \sum_{\ell \in \mathcal{L}_u} D[\ell, k, u] > 0\}, \quad |\mathcal{Q}_k| \leq U. \quad (65)$$

Conforme a abordagem de DCC empregada em sistemas CF-mMIMO escaláveis [116], a construção do *cluster* virtual é iniciada pela escolha de um *master AP*. Neste trabalho, o *master AP* do UE k é escolhido como o AP com maior ganho de larga escala:

$$\ell_k^* = \arg \max_{\ell} \beta_{\ell k}. \quad (66)$$

A CPU responsável pela coordenação dos recursos do UE k é a CPU que gerencia o seu *master AP*:

$$u_k^* = u \quad \text{tal que} \quad \ell_k^* \in \mathcal{L}_u. \quad (67)$$

Um mesmo AP pode atuar como *master AP* para mais de um UE, dependendo da distribuição espacial dos usuários e dos ganhos de canal. Além disso, mesmo quando não atua como *master*, um AP pode pertencer aos *clusters* virtuais de múltiplos UEs.

No contexto multi-CPU, a CPU u_k^* é responsável pela coordenação inicial dos recursos associados ao UE k . Quando o *cluster* virtual $\mathcal{V}_k$ contém APs pertencentes a outras CPUs, a alocação de recursos exige algum nível de coordenação entre as CPUs em $\mathcal{Q}_k$.

10.2.4 Transmissão em *downlink*

Na transmissão em *downlink*, o AP ℓ transmite o sinal

$$\mathbf{x}_\ell = \sum_{k=1}^K \mu[\ell, k] \hat{\mathbf{w}}_{\ell k} s_k \in \mathbb{C}^N, \quad (68)$$

em que $s_k \sim \mathcal{CN}(0, 1)$ é o símbolo de dado normalizado para o UE k ; $\hat{\mathbf{w}}_{\ell k} \in \mathbb{C}^N$ é o vetor de precodificação unitário ($\|\hat{\mathbf{w}}_{\ell k}\| = 1$) do AP ℓ para o UE k ; e $\mu[\ell, k] \geq 0$ é o coeficiente de potência. A potência total transmitida pelo AP ℓ é limitada por

$$\sum_{k=1}^K \mu[\ell, k]^2 = \|\boldsymbol{\mu}[\ell, :]\|_2^2 \leq p_d, \forall \ell.$$

O UE k recebe o sinal de todos os L APs:

$$y_k = \sum_{\ell=1}^L \mathbf{h}_{\ell k}^H \mathbf{x}_\ell + n_k = \sum_{\ell=1}^L \sum_{i=1}^K \mu[\ell, i] \underbrace{\mathbf{h}_{\ell k}^H \hat{\mathbf{w}}_{\ell i}}_{=G[\ell, k, i]} s_i + n_k, \quad (69)$$

em que $n_k \sim \mathcal{CN}(0, \sigma^2)$ é o ruído branco aditivo gaussiano, e define-se o tensor de fase instantânea $G[\ell, k, i] = \mathbf{h}_{\ell k}^H \hat{\mathbf{w}}_{\ell i} \in \mathbb{C}$, que representa o acoplamento do canal do UE k com o precodificador do UE i no AP ℓ .

O sinal recebido pode ser separado em três componentes distintas:

$$y_k = \underbrace{\text{DS}_k}_{\text{sinal desejado}} \cdot s_k + \underbrace{\text{BU}_k}_{\text{incerteza de beamforming}} \cdot s_k + \sum_{k' \neq k} \underbrace{\text{UI}_{k,k'}}_{\text{interferência inter-UE}} \cdot s_{k'} + n_k, \quad (70)$$

com cada componente definido como:

$$\text{DS}_k \triangleq \sum_{\ell \in \mathcal{V}_k} \mu[\ell, k] \mathbb{E}[G[\ell, k, k]], \quad (71)$$

$$\text{BU}_k \triangleq \sum_{\ell \in \mathcal{V}_k} \mu[\ell, k] \left(G[\ell, k, k] - \mathbb{E}[G[\ell, k, k]] \right), \quad (72)$$

$$\text{UI}_{k,k'} \triangleq \sum_{\ell \in \mathcal{V}_{k'}} \mu[\ell, k'] \mathbb{E}[G[\ell, k, k']]. \quad (73)$$

Com tais componentes, pode-se estabelecer a *Signal to Interference plus Noise Ratio* (SINR) estatística como:

$$\text{SINR}_k^{\text{stat}} = \frac{|\text{DS}_k|^2}{\mathbb{E}[|\text{BU}_k|^2] + \sum_{k' \neq k} \mathbb{E}[|\text{UI}_{k,k'}|^2] + \sigma^2}. \quad (74)$$

10.2.5 Problema de Otimização

O problema de alocação de potência com critério de justiça *max-min* é formulado como

$$\begin{aligned} & \underset{\boldsymbol{\mu} \in \mathbb{R}_+^{L \times K}}{\text{maximizar}} && \min_{k=1, \dots, K} \text{SE}_k(\boldsymbol{\mu}) \\ & \text{sujeito a} && \|\boldsymbol{\mu}[\ell, :]\|_2^2 \leq p_d, \quad \forall \ell, \\ & && \mu[\ell, k] = 0, \quad \text{se } D_{\text{flat}}[\ell, k] = 0. \end{aligned} \quad (75)$$

A eficiência espectral do UE k é dada por

$$\text{SE}_k(\boldsymbol{\mu}) = \frac{\tau_c - \tau_p}{\tau_c} \log_2(1 + \text{SINR}_k^{\text{stat}}(\boldsymbol{\mu})), \quad (76)$$

em que τ_c é o comprimento do bloco de coerência e τ_p é o comprimento da sequência de pilotos.

O problema (75) é não convexo, pois a função objetivo depende de razões quadráticas acopladas entre os coeficientes de potência de todos os usuários. Além disso, em uma arquitetura multi-CPU, cada CPU possui acesso direto apenas aos canais, APs e variáveis locais sob sua responsabilidade. Portanto, uma solução centralizada exigiria o compartilhamento de informações globais de canal, associação e potência entre CPUs, aumentando significativamente a carga de sinalização no *fronthaul* e no *backhaul*.

Por esse motivo, um dos objetivos do trabalho é a proposição de um método que aproxime a solução de (75) com baixa complexidade computacional e baixo *overhead* de comunicação entre CPUs, preservando a estrutura *user-centric* definida por D_{flat} e respeitando as restrições de potência por AP.

10.3 Métodos Propostos

Esta seção descreve os três métodos avaliados neste trabalho para reduzir a complexidade de processamento e o *overhead* de comunicação em sistemas CF-mMIMO com múltiplas CPUs. O primeiro método consiste em uma estratégia escalável de DCC, na qual cada UE é servido apenas por um subconjunto limitado de APs. O segundo método estende essa estratégia ao considerar explicitamente a partição dos APs entre CPUs, buscando reduzir a quantidade de CPUs envolvidas no atendimento de cada UE. Por fim, o terceiro método trata da alocação de potência distribuída, com foco no critério de justiça *max-min* e em baixa troca de informação entre CPUs.

10.3.1 DCC Escalável

O método proposto tem por objetivo contruir, para cada UE, um *cluster* virtual de APs que seja simultaneamente eficiente em termos de qualidade de canal e compatível com a estrutura multi-CPU da rede. Diferentemente do critério convencional, baseado apenas nos maiores valores de desvanecimento em larga escala, o método proposto combina informação de canal estatístico e informação geométrica para induzir associações AP-UE mais alinhadas aos domínios de CPU.

Para cada par (ℓ, k) , define-se o vetor de características

$$\mathbf{s}_{\ell,k} = ((1 - \alpha)\beta_{\ell k}, \alpha d_{\ell k}, \alpha(1 - \alpha)\beta_{\ell k}d_{\ell k}) \in \mathbb{R}^3, \quad (77)$$

em que $\alpha \in [0, 1]$ controla o peso relativo da distância em relação ao $\beta_{\ell k}$. A terceira componente, $\beta_{\ell k}d_{\ell k}$, introduz uma interação não linear entre qualidade de canal e proximidade geométrica, permitindo ao algoritmo capturar padrões que não seriam observados por uma ordenação direta baseada em apenas uma dessas métricas.

A escolha dessas características é motivada pela restrição de baixa cooperação entre CPUs. Métricas baseadas em taxa instantânea ou em CSI global exigiriam troca de informação entre CPUs, o que contraria o objetivo de escalabilidade. Em contraste, $\beta_{\ell k}$ é uma grandeza de larga escala, estimada localmente a partir dos pilotos de *uplink*, enquanto $d_{\ell k}$ pode ser obtida a partir dos metadados de posicionamento da rede. Assim, o vetor $\mathbf{s}_{\ell,k}$ pode ser computado sem compartilhamento de CSI instantânea entre CPUs.

Para cada UE k , constrói-se a matriz de características

$$\mathbf{X}_k = \begin{bmatrix} \mathbf{s}_{1,k}^T \\ \mathbf{s}_{2,k}^T \\ \vdots \\ \mathbf{s}_{L,k}^T \end{bmatrix} \in \mathbb{R}^{L \times 3}. \quad (78)$$

Em seguida, aplica-se o algoritmo *K-means* a $\mathbf{X}_k$, particionando os L APs em G grupos. O *master AP* do UE k é definido como o AP com maior LSFC:

$$\ell_k^* = \arg \max_{\ell} \beta_{\ell k}. \quad (79)$$

Seja g_k^* o rótulo do grupo *K-means* ao qual pertence ℓ_k^* . O conjunto candidato de APs para o UE k é então dado por

$$\mathcal{C}_k = \{\ell : c_{\ell,k} = g_k^*\}, \quad (80)$$

em que $c_{\ell,k}$ é o rótulo atribuído ao AP ℓ no agrupamento realizado para o UE k .

O *cluster* virtual $\mathcal{V}_k$ é formado selecionando, dentro de $\mathcal{C}_k$, os L_{cap} APs mais próximos do centroide $\mu_{g_k^*}$ do grupo do *master AP*:

$$\mathcal{V}_k = \text{Top}_{L_{\text{cap}}} \left\{ -\|\mathbf{s}_{\ell,k} - \mu_{g_k^*}\|_2 : \ell \in \mathcal{C}_k \right\}. \quad (81)$$

A notação acima indica que são escolhidos os L_{cap} APs com menor distância ao centroide do grupo selecionado. Por fim, define-se $D[\ell,k,u]$ como

$$D[\ell,k,u] = \begin{cases} 1, & \ell \in \mathcal{V}_k \text{ e } \ell \in \mathcal{L}_u, \\ 0, & \text{caso contrário.} \end{cases} \quad (82)$$

10.3.2 DCC Escalável com CPU-aware

A versão CPU-aware do DCC escalável estende o vetor de características usado no agrupamento não supervisionado, incorporando uma componente explícita associada ao domínio de CPU. O objetivo é favorecer a seleção de APs pertencentes à mesma CPU do *master AP* do UE, reduzindo a quantidade de CPUs envolvidas no atendimento de um mesmo usuário e, consequentemente, o *overhead* das etapas posteriores de coordenação.

Em relação à versão anterior, concatenam-se as componentes $c_{\ell k}$ e $t_{\ell k}$ definidas como:

$$c_{\ell k} = \text{zscore}_{\ell} \left(\mathbf{1}[\ell \in \mathcal{L}_{u^*(k)}] \right), \quad (83)$$

$$t_{\ell k} = \text{zscore}_{\ell} \left(\beta_{\ell k} - \max_{\ell'} \beta_{\ell' k} \right), \quad (84)$$

em que $c_{\ell k}$ indica se o AP ℓ pertence ao domínio da CPU responsável pelo *master AP* do UE k , normalizado para variância unitária ao longo de ℓ . O termo $u^*(k)$ é o índice da CPU tal que $\ell_k^* \in \mathcal{L}_{u^*(k)}$, ou seja, a CPU responsável pelo *master AP* do UE k .

A característica $t_{\ell k}$ mede a proximidade de cada AP ao topo do ranking de β , diferenciando regiões de menor e maior ganho.

O vetor de características submetido ao K-means é, então

$$\mathbf{x}_{\ell k} = [\mathbf{s}_{\ell k}, \gamma_{\text{cpu}} c_{\ell k}, \tau_{\text{top}} t_{\ell k}], \quad (85)$$

em que $\gamma_{\text{cpu}} \geq 0$ regula a atração por APs co-gerenciados pelo mesmo CPU do *master AP*, e $\tau_{\text{top}} \geq 0$ reforça a separação dos APs de maior β no espaço de características.

Procedimento semelhante ao descrito anteriormente é aplicado aqui para a seleção dos melhores APs.

10.3.3 Alocação de Potência Distribuída

A estratégia de α -refinement é utilizada como uma etapa de refinamento de baixa dimensão para melhorar a justiça da alocação de potência em *downlink*. Em vez de otimizar diretamente todos os coeficientes $\mu[\ell,k]$, o método parte de uma solução base factível $\mu^{(0)} \in \mathbb{R}_+^{L \times K}$ e ajusta apenas um vetor de escala por usuário,

$$\boldsymbol{\alpha} = [\alpha_1, \dots, \alpha_K]^T \in \mathbb{R}_+^K. \quad (86)$$

Assim, a alocação refinada é parametrizada por

$$\mu_{\ell k}(\boldsymbol{\alpha}) = \mu_{\ell k}^{(0)} \alpha_k, \quad \forall \ell, k. \quad (87)$$

Essa parametrização preserva a estrutura de *cluster* definida por D_{flat} , pois se $\mu_{\ell k}^{(0)} = 0$ para um par AP-UE, então $\mu_{\ell k}(\boldsymbol{\alpha}) = 0$ para qualquer valor de α_k .

A ideia central é favorecer os usuários com menor eficiência espectral corrente. Para uma dada iteração t , calcula-se

$$\text{SE}_k^{(t)} = \text{SE}_k(\boldsymbol{\mu}^{(0)} \text{diag}(\boldsymbol{\alpha}^{(t)})), \quad (88)$$

e define-se o peso relativo

$$w_k^{(t)} = \frac{1/\text{SE}_k^{(t)}}{\frac{1}{K} \sum_{j=1}^K 1/\text{SE}_j^{(t)}}. \quad (89)$$

Dessa forma, usuários com menor SE_k recebem pesos maiores, enquanto usuários com maior SE_k recebem pesos menores. O vetor $\boldsymbol{\alpha}$ é então atualizado por uma regra multiplicativa com amortecimento:

$$\alpha_k^{(t)} \leftarrow \alpha_k^{(t)} \left[(1 - \eta) + \eta w_k^{(t)} \right], \quad 0 < \eta \leq 1, \quad (90)$$

em que η controla a intensidade da atualização.

Após a atualização, é necessário garantir que a restrição de potência por AP, $P_\ell(\cdot)$, continue sendo satisfeita. Após isso, para preservar a factibilidade, aplica-se uma projeção escalar global

$$\text{proj}^{(t+1)} = \min \left\{ 1, \min_{\ell: P_\ell(\boldsymbol{\alpha}^{(t)}) > 0} \sqrt{\frac{p_d}{P_\ell(\boldsymbol{\alpha}^{(t)})}} \right\}, \quad (91)$$

e atualiza-se

$$\alpha_k^{(t+1)} \leftarrow \text{proj}^{(t+1)} \alpha_k^{(t)}. \quad (92)$$

Com isso, para todo AP ℓ ,

$$\sum_{k=1}^K \left(\mu_{\ell k}^{(0)} \alpha_k^{(t+1)} \right)^2 \leq p_d. \quad (93)$$

10.4 Resultados

Nesta seção, são apresentados resultados parciais obtidos durante o desenvolvimento das estratégias escaláveis de formação de *clusters* e alocação de potência. As simulações consideram um sistema CF-mMIMO multi-CPU com os parâmetros listados na Tabela 4.

Tabela 4: Parâmetros de simulação.

Parâmetro	Símbolo	Valor
Número de APs	L	100
Antenas por AP	N	4
Número de UEs	K	20
Número de CPUs	U	4
Lado da área de cobertura	—	2500 m
Potência por AP	p_d	100 mW
Comprimento do bloco de coerência	τ_c	200 amostras
Comprimento dos pilotos	τ_p	5 ou 20 amostras
Amostras de Monte Carlo	—	1.000

A primeira análise compara o método *baseline*, baseado apenas nos coeficientes de larga escala $\beta_{\ell k}$, com o método proposto de DCC escalável, que combina informação de canal e distância AP-UE por meio do agrupamento *K-means*. A Figura 26 apresenta a SE média e o índice de justiça de Jain (JFI) em função do número de APs no *cluster*, L_{cap} . Para avaliar a sensibilidade da formação de *clusters* à técnica adotada na transmissão, os resultados apresentados também avaliam os precodificadores *Maximum Ratio* (MR) e *Regularized Zero-Forcing* (RZF), permitindo comparar a proposta com soluções de precodificação de diferentes complexidades [120]. Observa-se que o método proposto obtém desempenho superior ao *baseline*, indicando que a combinação entre $\beta_{\ell k}$ e distância produz *clusters* mais adequados à operação distribuída. Esses resultados correspondem ao cenário com pilotos ortogonais, isto é, sem contaminação de pilotos.

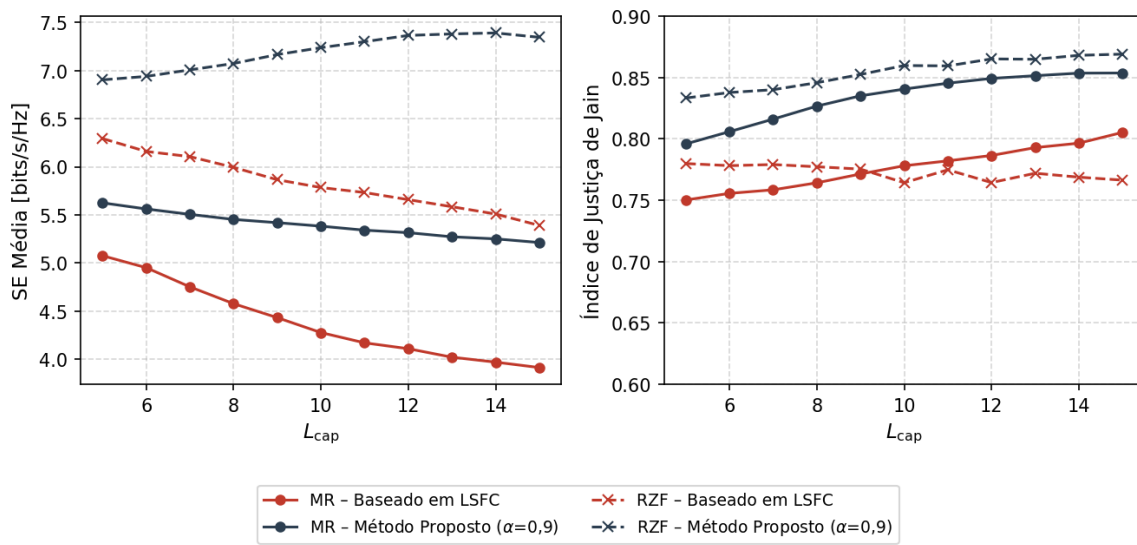


Figura 26: Comparação de SE e JFI para o método de DCC escalável.

A contaminação de pilotos faz com que a estimação de canal seja prejudicada e a análise apenas por β fique prejudicada. Por isso, a Tabela 5 mostra a comparação entre o método de DCC escalável, DCC CPU-aware e o *baseline*. Observa-se que a adição das características γ_{cpu} e τ_{top} fazem diferença em cenários onde a estimação de canal é comprometida.

Por fim, podemos observar na Figura 27, o método α -refinement para o processo de alocação de potência. Considera-se aqui que as CPUs não se comunicam entre si, *a priori*. A primeira estimativa da alocação de potências para cada UE é feita individualmente por cada CPU. Como um UE pode ser servido por APs de mais de uma CPU, a alocação pode não ser única. Então, como proposta de alocação de potências com o mínimo de *overhead* se propõe aqui a média entre todas as alocações realizadas. Na Figura 27, a curva CPU-Avg é a alocação por este método. Neste caso, em cada CPU, realiza-se o algoritmo de *bisection* para a maximização da SE. As demais curvas do gráfico demonstram as iterações do α -refinement. O passar das iterações aumenta o *overhead* da operação mas o resultado é incremental.

10.5 Conclusão

Este trabalho apresentou uma opção escalável para formação de *clusters* e alocação de potência em sistemas CF-mMIMO com múltiplas CPUs. O método de DCC escalável utiliza aprendizado não supervisionado para construir *clusters* virtuais a partir de informações de larga

Tabela 5: Comparação de desempenho usando precodificador MR e $\tau_p = 5$ (com contaminação de pilotos). Parâmetro DCC escalável $\alpha = 0,9$. Parâmetros do DCC CPU-*aware*: $\gamma_{\text{cpu}} = 0.1$, $\tau_{\text{top}} = 0.3$.

L_{cap}	Método	Média. SE	Mínimo . SE	JFI
5	DCC (K-Means)	1,3635	0,1234	0,7651
	LSFC- <i>based</i>	1,2309	0,0710	0,7360
	DCC CPU- <i>aware</i>	1,3647	0,1245	0,7656
7	DCC (K-Means)	1,3996	0,1394	0,7700
	LSFC- <i>based</i>	1,2434	0,0739	0,7364
	DCC CPU- <i>aware</i>	1,4004	0,1412	0,7705
10	DCC (K-Means)	1,4181	0,1529	0,7714
	LSFC- <i>based</i>	1,2473	0,0750	0,7354
	DCC CPU- <i>aware</i>	1,4197	0,1541	0,7722
15	DCC (K-Means)	1,4265	0,1568	0,7730
	LSFC- <i>based</i>	1,2429	0,0756	0,7336
	DCC CPU- <i>aware</i>	1,4281	0,1578	0,7736

escala, combinando ganho de canal, distância AP-UE e, na versão CPU-*aware*, uma componente associada ao domínio de CPU. Com isso, busca-se preservar os ganhos de macrodiversidade do CF-mMIMO, ao mesmo tempo em que se reduz a necessidade de coordenação entre diferentes CPUs.

Também foi descrita uma estratégia de alocação de potência baseada em α -*refinement*, na qual uma solução inicial factível é refinada por meio de fatores escalares por usuário. Essa abordagem reduz a dimensão do problema de alocação de potência e favorece usuários com menor eficiência espectral, tornando a solução mais adequada ao critério de justiça *max-min*.

De forma geral, os métodos apresentados apontam para uma implementação mais prática de redes CF-mMIMO distribuídas, pois limitam o compartilhamento de informações globais, exploram apenas LSF na etapa de *clustering* e mantêm a alocação de potência em uma forma computacionalmente simples. Como continuidade, pretende-se aprofundar a análise do compromisso entre eficiência espectral, justiça entre usuários e *overhead* de sinalização em cenários com maior densidade de APs, reutilização de pilotos e diferentes níveis de cooperação entre CPUs.

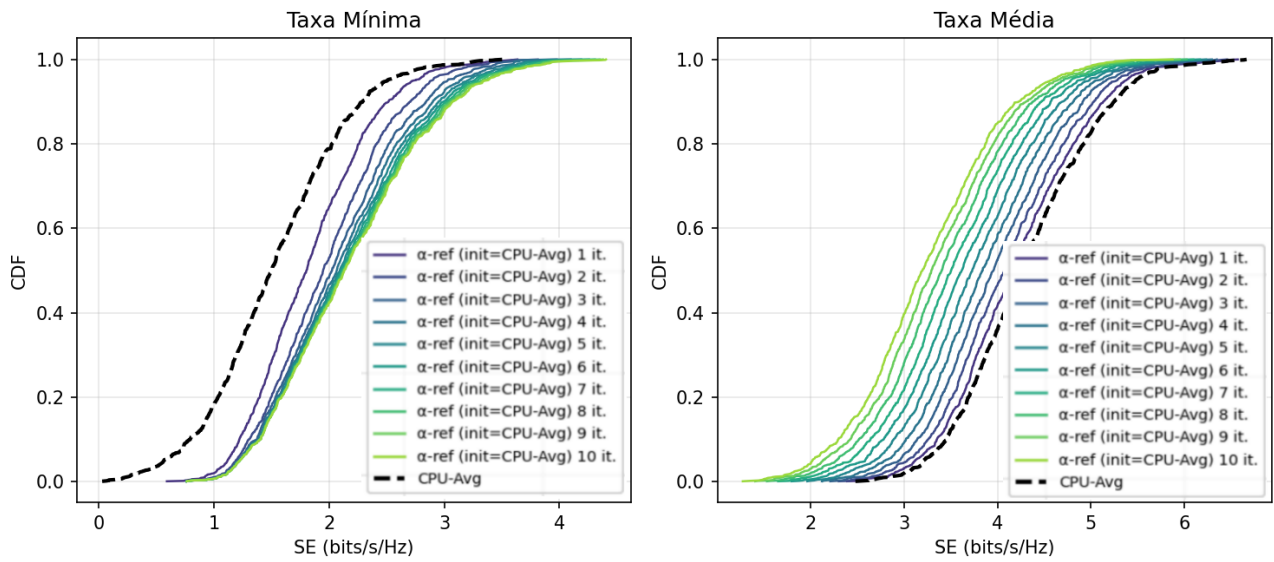


Figura 27: Alocação de potência segundo o método proposto e o método denominado como CPU-Avg.

11 Análise de Generalização e Divergência de Domínio em Seleção de Feixe Baseada em Informações de Contexto

Joanna Manjarres, José Ferreira de Rezende
UFRJ

11.1 Introdução

A evolução das comunicações móveis para a sexta geração (6G) estabelece requisitos extremamente rigorosos de desempenho, incluindo taxas de transmissão da ordem de Tbps e latências ultrabaixas inferiores a 0,1 ms. Nesse cenário, a utilização de bandas de ondas milimétricas (mmWave) surge como uma solução promissora para atender às demandas de capacidade e densidade das futuras redes sem fio. Entretanto, as características de propagação em altas frequências introduzem desafios significativos, uma vez que os sinais mmWave apresentam elevada sensibilidade a bloqueios, reflexões, mobilidade dos usuários e variações ambientais, comprometendo a estabilidade e confiabilidade dos enlaces de comunicação. Diante dessas limitações, técnicas de formação de feixe (beamforming) tornam-se essenciais para concentrar energia na direção de transmissão desejada e compensar as elevadas perdas de propagação. Como consequência, a seleção eficiente do feixe passa a desempenhar um papel central na manutenção da conectividade, especialmente em cenários veiculares dinâmicos, nos quais mudanças rápidas do ambiente exigem mecanismos de adaptação contínua e baixa latência.

Durante a fase anterior do projeto Brasil 6G, foi conduzida uma investigação voltada à avaliação da robustez e da capacidade de generalização de diferentes modelos de aprendizado de máquina aplicados à seleção de feixe utilizando informações de contexto, como coordenadas de posicionamento e dados LiDAR. Foram analisadas abordagens baseadas principalmente em arquiteturas deep learning, utilizadas como baseline a partir dos trabalhos propostos por Saleh Batool [3] e Ruseckas [4], além da rede neural sem peso WiSARD, proposta desenvolvida ao longo desta pesquisa. As avaliações experimentais foram realizadas utilizando dados sintéticos gerados pelo dataset Raymobtime, permitindo investigar o comportamento dos modelos sob diferentes condições de propagação e mudanças de domínio entre cenários distintos.

Os resultados mostraram que:

- entradas baseadas apenas em coordenadas apresentam elevada sensibilidade à troca de domínios;
- dados LiDAR proporcionam maior estabilidade entre cenários;
- a combinação multimodal (LiDAR + coordenadas) oferece os melhores resultados de generalização;
- todos os modelos sofrem degradação significativa em avaliações cruzadas entre cenários, principalmente em condições LoS.

Esses resultados evidenciaram que a limitação de generalização não está associada exclusivamente ao modelo utilizado, mas também às propriedades estatísticas e estruturais dos próprios datasets. Dessa forma, surgiu a necessidade de desenvolver uma abordagem centrada em dados (data-centric analysis) capaz de investigar como características distribucionais, geométricas e estruturais dos conjuntos de dados afetam o desempenho dos modelos de aprendizado de máquina. Assim, durante esta fase, foi conduzida uma investigação aprofundada sobre as

propriedades dos datasets s008 e s009 do Raymobtime, considerando diferentes condições de propagação (LoS, NLoS e ALL), com foco na caracterização do espaço de feixes, análise de divergência entre distribuições e estudo de métricas capazes de prever dificuldades de generalização em cenários 6G.

11.2 Desenvolvimento das Atividades

Com o objetivo de investigar, de forma sistemática e estatisticamente fundamentada, o impacto das divergências entre domínios no desempenho da seleção de feixes, foi proposta uma metodologia baseada em um protocolo de avaliação intercenários. Essa abordagem permite caracterizar tanto as diferenças estatísticas entre os datasets quanto sua influência sobre a capacidade de transferência do modelo, estabelecendo uma relação entre métricas de similaridade de dados e o desempenho obtido em avaliações cruzadas.

Para formalizar as mudanças de distribuição, definimos um domínio $\mathcal{D} = \{\mathcal{X}, P(X)\}$, em que $\mathcal{X}$ é o espaço de características e $P(X)$ é a distribuição de probabilidade marginal das entradas contextuais. A tarefa de aprendizado $\mathcal{T} = \{\mathcal{Y}, P(Y|X)\}$ visa capturar a probabilidade condicional do feixe ótimo dado o contexto observado. Nessa análise, definimos o cenário s008 como Domínio de Origem ($\mathcal{D}_S$) e o cenário s009 como Domínio de Destino ($\mathcal{D}_T$). Para avaliar a robustez do modelo além da validação localizada padrão, nossa metodologia evita divisões aleatórias arbitrárias e executa um protocolo de três níveis que combina treinamento estratificado e testes entre cenários:

Nível 1 — *Baseline* de Desempenho: Os modelos são treinados utilizando o conjunto de treinamento completo e não filtrado do Domínio de Origem $\mathcal{D}_S = s008$ e avaliados no conjunto de teste completo do Domínio de Destino $\mathcal{D}_T = s009$. Essa configuração estabelece o referencial clássico de generalização, no qual uma única rede deve lidar simultaneamente com diversas misturas de propagação.

Nível 2 — *Dinâmica de Propagação:* Os conjuntos de dados de origem e destino são fisicamente particionados em subdistribuições puras de LoS (visada direta) e NLoS (sem visada direta). Redes especializadas são então treinadas exclusivamente com os componentes LoS de $\mathcal{D}_S$ e avaliadas nos componentes LoS de $\mathcal{D}_T$ (com um fluxo de processamento isolado idêntico implementado para NLoS). Esse nível revela o limite máximo de desempenho quando os modelos são otimizados em meios de propagação estruturalmente homogêneos.

Nível 3 — *Análise de Generalização (Troca de Domínios):* Para expor vulnerabilidades latentes dos dados e avaliar a adaptação bidirecional, realizamos uma troca abrangente de domínios. Aqui, o cenário denso e com obstruções severas s009 é designado como o Domínio de Origem ($\mathcal{D}_S$) para a otimização do modelo, enquanto o cenário equilibrado s008 atua como o Domínio de Destino ($\mathcal{D}_T$) não visto anteriormente. Essa inversão é aplicada simetricamente às configurações ALL (todos os dados), LoS e NLoS, testando se um grande volume de amostras se traduz em resiliência entre cenários distintos.

11.2.1 Configuração experimental multimodal

Este estudo utiliza a estrutura do conjunto de dados Raymobtime¹, uma plataforma de simulação abrangente projetada para pesquisas em comunicação veicular 5G/6G [121], amplamente usado na literatura. O Raymobtime integra padrões realistas de mobilidade veicular

¹<https://www.lasse.ufpa.br/raymobtime/>

com a modelagem determinística de canal baseada em *ray-tracing* para gerar dados multimodais sincronizados. Uma vantagem fundamental dessa estrutura é sua capacidade de garantir realizações de canal consistentes ao longo do tempo, da frequência e do espaço, tornando-a particularmente adequada para avaliar a capacidade de generalização *out-of-distribution* (OOD) de modelos de seleção de feixes sob variações ambientais não estacionárias. Uma descrição detalhada da metodologia de *ray-tracing* subjacente e das estruturas de dados baseadas em HDF5 pode ser encontrada em [122].

O *dataset* Raymobtime contém 9 cenários distintos, identificados como s000 a s009. Para este trabalho, foram selecionados os subconjuntos s008 e s009, nos quais os dados foram coletados em um ambiente que simula uma rua tipo cânion, localizada em Rosslyn, Virginia. Este conjunto de dados é composto por informações de traçado de raios, nuvem de pontos providas por sensores LiDAR, imagens de câmeras e informações de posição (x,y,z) dos receptores. Por tratar-se de uma simulação realista, nestes subconjuntos se encontram situações de comunicação dos tipos LoS e NLoS entre um par Tx-Rx. Na Tabela 6 estão resumidas as estatísticas gerais dos *datasets* em relação às amostras LoS e NLoS.

Dataset	# Amostras	Canais LOS	Canais NLOS
s008	11194	6482	4712
s009	9638	1473	8165

Tabela 6: Estatísticas dos *dataset* s008 e s009.

11.2.2 Arquiteturas de *baseline* Multimodais Avaliadas

Selecionamos três modelos competitivos e estruturalmente diversos da literatura para avaliar o protocolo entre domínios, representando diferentes paradigmas de aprendizado para a previsão de feixes sensível ao contexto:

Framework ResNet-MLP (B. Salehi et al. [3]): Emprega uma arquitetura profunda de múltiplos ramos, na qual uma espinha dorsal (backbone) ResNet-18 extrai padrões espaciais de alto nível a partir de grades LiDAR voxelizadas, enquanto um ramo paralelo de MLP densa processa as coordenadas GNSS. As representações latentes resultantes são concatenadas para a classificação final em 256 classes.

Rede Convolutacional de Ramos Duplos (J. Ruseckas et al. [4]): Utiliza uma topologia assimétrica de dois ramos. Uma MLP leve e com poucos parâmetros mapeia a confiabilidade das coordenadas, enquanto uma CNN 3D densa e personalizada extrai restrições geométricas das nuvens de pontos LiDAR. As representações são fundidas antes da inferência na camada de saída.

Framework de Rede Neural sem Pesos (J. Manjarres et al. [5]): Diferindo da otimização por descida de gradiente, este paradigma utiliza Redes Neurais sem Pesos (WNNs) baseadas em RAM. Ele mapeia padrões de entrada diretamente em tabelas de consulta (lookup tables) binárias dentro de nós de memória de hardware. Essa abordagem centrada na memória elimina a retropropagação, proporcionando uma inferência determinística e de baixa complexidade, adequada para nós V2I 6G em tempo real.

As configurações críticas de hiperparâmetros, incluindo configurações de lote, taxas de aprendizado e definições de mapeamento de memória — estão resumidas na Tabela 7, garantindo total transparência experimental e reprodutibilidade.

Parâmetros	B. Salehi	J. Ruseckas	J. Manjarres
Épocas de treinamento	100	50	-
Tamanho do lote	32	32	-
Taxa de aprendizado	0,0001	Increase according to a coise law for the 30% of learning duration	-
Otimizador	Adam	Adam	-
	β (0.9 , 0.999)	β (min 0.85)	-
Função de ativação na saída	Categorical cross entropy		-
Tamanho do endereço da memória	-	-	32 bits
Resolução do termômetro	-	-	8 bits
Tamanho do modelo	18,4 MB	4,25 MB	6551 bits

Tabela 7: Parâmetros usados pelas referencias B. Salehi [3], J. Ruseckas [4] and J. Manjarres[5] nos modelos multimodais

11.2.3 Avaliação Empírica e a Barreira da Generalização entre Domínios

Ao transitar da avaliação padrão para meios de propagação isolados e trocas de domínio bidirecionais, mapeamos os limites de desempenho da seleção de feixes sensível ao contexto e evidenciamos as limitações severas da avaliação de conjuntos de dados centrada no volume.

- **Resultados da etapa 1: *Baseline* de desempenho em condições gerais** Na etapa inicial, os modelos foram avaliados sob uma estrutura de referência *baseline* tradicional, na qual o treinamento é realizado no conjunto de treinamento completo e não filtrado do domínio de origem ($\mathcal{D}_S = s008$) e a avaliação ocorre no conjunto de teste correspondente do domínio de destino ($\mathcal{D}_T = s009$). Essa configuração representa um cenário clássico de implementação, no qual uma única arquitetura neural deve resolver simultaneamente o mapeamento de feixes em condições de enlaces de propagação mistos e não separados. A Tabela 8 resume a precisão geral de predição e as métricas de tempo computacional de treinamento para cada abordagem.

Model	Accuracy			Throughput Ratio			Training time	Total params
	Top-1	Top-5	Top-10	Top-1	Top-5	Top-10		
B. Salehi [3]	42,6%	74,7%	82,7%	0,56	0,82	0,87	2301,6	18,4 MB
J. Ruseckas [4]	58,8%	87,2%	92,5%	0,73	0,92	0,95	2180,0	4,25 MB
J. Manjarres [5]	60,0%	87,1%	92,1%	0,74	0,92	0,95	2,9	6551 bits

Tabela 8: Comparação entre as técnicas do estado da arte para métricas-chave de entrada multimodal

- **Resultados da Etapa 2: Impacto da Dinâmica de Propagação do Canal**

Para compreender como as características físicas do enlace sem fio afetam as previsões contextuais, a Etapa 2 divide ambos os cenários em distribuições explícitas e isoladas de Visada Direta (LoS) e de Ausência de Visada Direta (NLoS). Modelos especializados foram otimizados e validados sob regimes de propagação estruturalmente homogêneos

(isto é, treinados em $s008_{LoS}$ e testados em $s009_{LoS}$, e de forma simétrica para NLoS). A Fig. 28 ilustra as métricas de acurácia Top-1 resultantes.

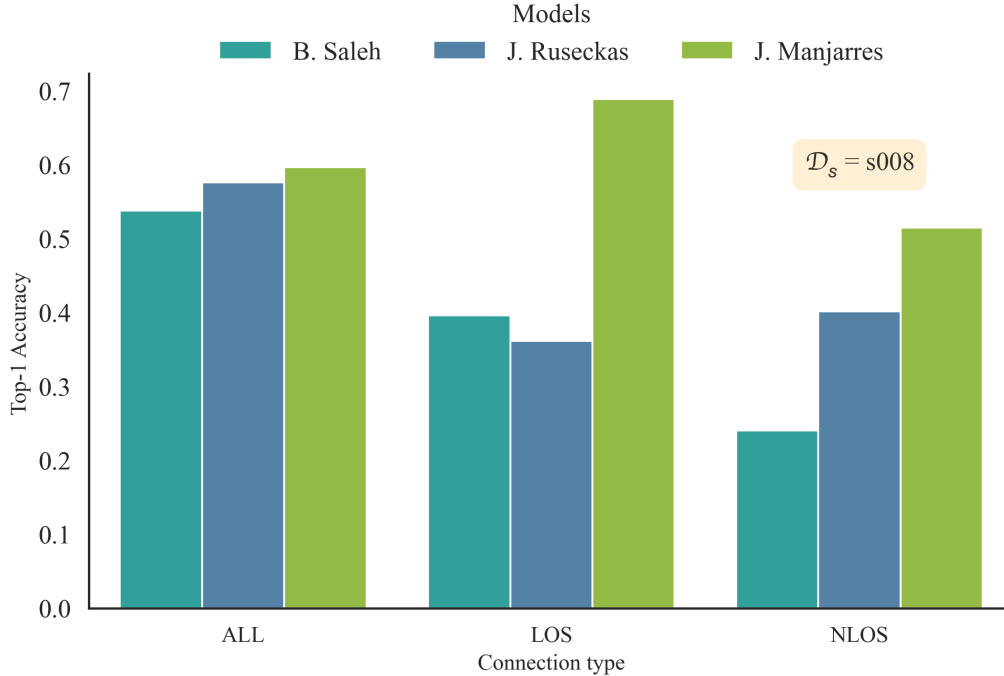


Figura 28: Performance comparison of ML models by isolated propagation type using multimodal context (LiDAR + GPS) with Domain Source $\mathcal{D}_S = s008$ and Domain Target $\mathcal{D}_T = s009$.

A análise isolada confirma que a física de propagação determina os limites de desempenho de características multimodais. Links sem obstruções reduzem a complexidade do mapeamento; o paradigma de rede neural sem pesos proposto por J. Manjarres tira proveito disso, alcançando uma precisão máxima de 68,90% no subdomínio LoS, enquanto cai para 51,47% em condições NLoS.

Além disso, os resultados destacam o benefício sinérgico da fusão de dados multimodais em situações de bloqueio severo de sinal. Por exemplo, no cenário de rastreamento NLoS, a suplementação de voxels geométricos de LiDAR com coordenadas GNSS, dentro da arquitetura de J. Ruseckas, melhorou a precisão da previsão de 36,00% para 40,17%. Essa mudança de comportamento indica que, quando bloqueios físicos eliminam pistas geométricas dominantes da nuvem de pontos LiDAR, a rede aprende a ponderar coordenadas espaciais em macroescala para restringir o espaço de busca do codebook de feixes viável.

• Resultados do Nível 3: Generalização entre Cenários e o Protocolo de Troca de Domínio

Para testar a resiliência do modelo contra deslocamentos espaciais do tipo *out-of-distribution* (OOD), o Nível 3 introduz um teste de estresse bidirecional ao executar uma troca de domínio abrangente, designando o cenário $s009$, estruturalmente denso e com muitas obstruções, como o domínio de treinamento de origem ($\mathcal{D}_S = s009$) e o cenário $s008$ como o alvo de avaliação inédito ($\mathcal{D}_T = s008$). Essa inversão revela vulnerabilidades assimétricas profundas em todos os modelos.

Conforme ilustrado na Fig. 29, quando a troca de domínio é aplicada aos conjuntos de dados não filtrados (ALL), ocorre uma degradação de desempenho acentuada e transversal nos três paradigmas em comparação com o Nível 1. No entanto, a vulnerabilidade mais crítica surge nos canais de propagação isolados. Em condições de propagação NLoS, todos os modelos avaliados sofrem um colapso catastrófico de desempenho, com as acurácias estagnando severamente em uma faixa estreita de 24% a 25%. Esse colapso ocorre apesar de o domínio de treinamento ($\mathcal{D}_S = s009$) conter uma abundância de dados NLoS (8.164 amostras), volume quase duas vezes maior do que o disponível no domínio alvo.

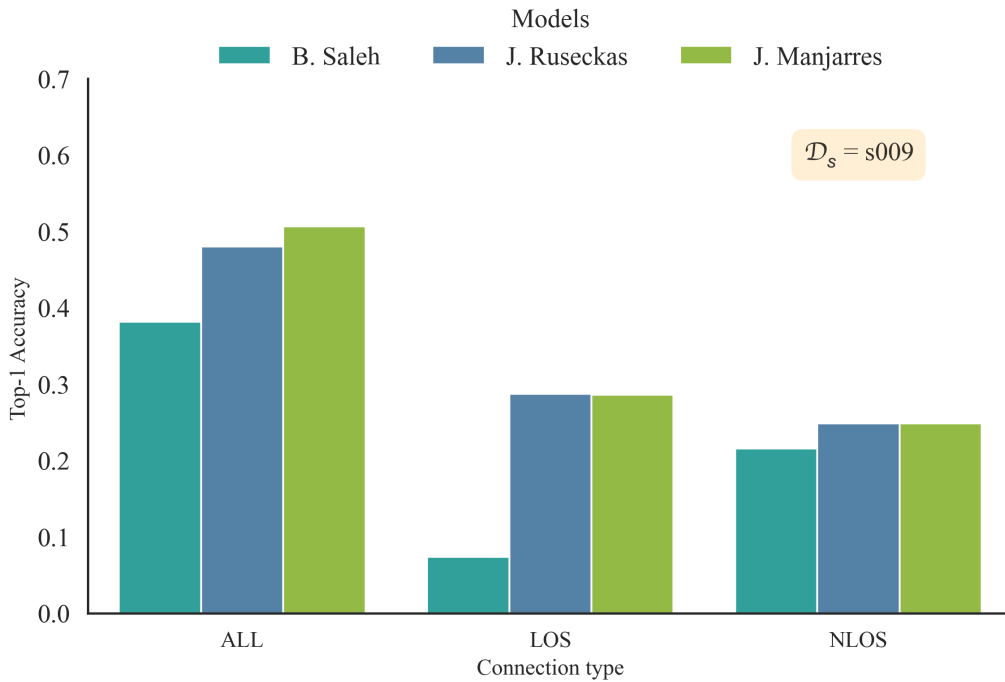


Figura 29: Comparação de desempenho entre domínios sob o protocolo de troca de domínios *Domain Swap*, utilizando contexto multimodal (LiDAR + GPS), com domínio de origem $\mathcal{D}_S = s009$ e domínio de destino $\mathcal{D}_T = s008$.

De forma simétrica, sob condições de visibilidade direta (LoS), atinge-se um limite estrito de desempenho. As arquiteturas profundas multimodais de J. Manjarres e J. Ruseckas convergem para limites de acurácia praticamente idênticos (28,60% e 28,72%, respectivamente). Essa convergência de comportamento demonstra que o gargalo de generalização durante uma transição de domínio não depende mais da profundidade da rede neural, da complexidade estrutural ou dos algoritmos de otimização, sendo, em vez disso, determinado por propriedades externas.

11.2.4 Caracterização Estatística dos Datasets

A primeira etapa da investigação concentrou-se na caracterização estatística das distribuições dos índices de feixe presentes nos datasets avaliados. Inicialmente, foi utilizada a entropia de Shannon normalizada para quantificar a incerteza associada às distribuições de feixes em diferentes cenários de propagação, seguidamente foi realizada uma avaliação de sobreposição de classes. Com o intuito de quantificar o desalinhamento das classes entre os dataset, foi usada

a divergência KL e a distância wassertein. Além de apenas avaliar as classes, foram usadas as métricas Optimal Transport Dataset Distance (OTDD) e Nearest Neighbor (kNN-Class)

Nas figuras 30 e 31 são apresentados os histogramas dos datasets *s008* e *s009* completos e filtrados por tipo de conexão. Neles, é possível observar que existe um forte desbalanceamento de dados com uma concentração maior na faixa de índice de feixe de 145 a 160, juntamente com classes raras e uma estrutura multimodal. No entanto, o cenário *s009* apresenta uma concentração mais acentuada, no qual o feixe dominante representa 18% das amostras em comparação com 12% em *s008*, sugerindo um espaço de rótulos menos diverso no domínio alvo. Também é evidente a desproporção de quantidade de amostras entre os dados LoS e NLoS para o dataset *s009*.

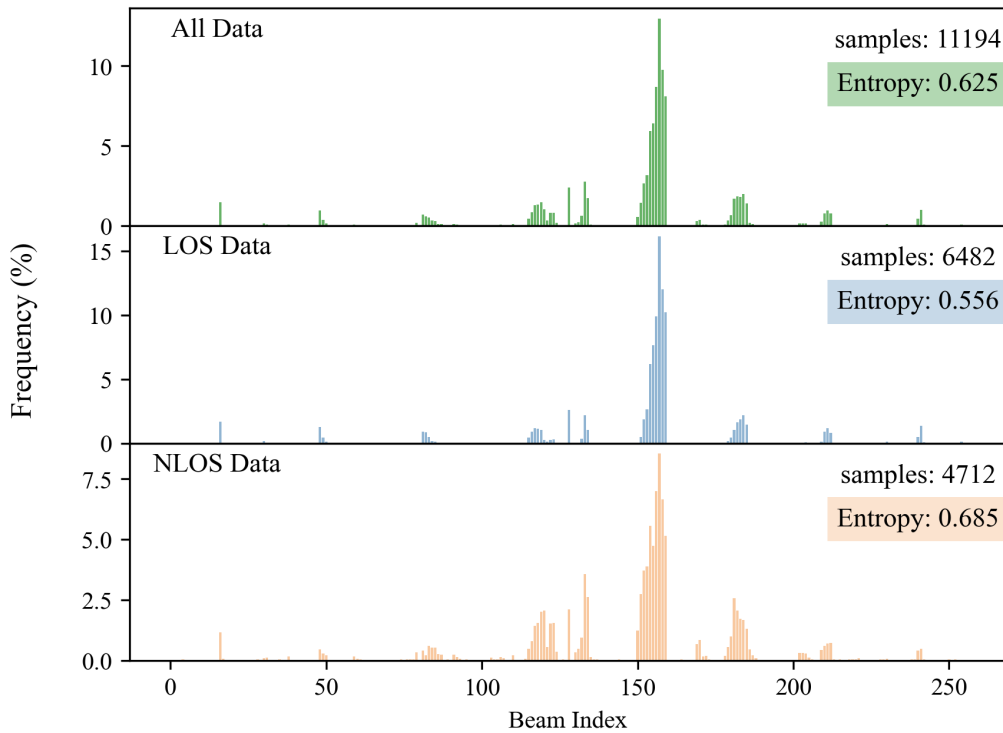


Figura 30: Distribuição de classes dos índices de feixe para o conjunto de dados *s008*.

Para quantificar essa diversidade, definimos a entropia de Shannon $H(X)$ para uma variável aleatória discreta X com probabilidades $P = \{p_1, \dots, p_n\}$ como:

$$H(X) = - \sum_{i=1}^n p_i \log_2 p_i \quad (94)$$

Neste estudo, utilizamos a entropia normalizada, $H_{\text{norm}} = H(X) / \log_2(n_{\text{classes}})$,

Essa normalização resulta em um valor entre 0 e 1, onde 1 representa a incerteza máxima (distribuição uniforme) e 0 reflete uma distribuição altamente previsível dominada por uma única classe.

Interpretamos a diferença de entropia ($|\Delta H|$) entre os conjuntos de treinamento e teste como uma aproximação da mudança estrutural:

- $|\Delta H| < 0,05$: Variação estrutural menor;

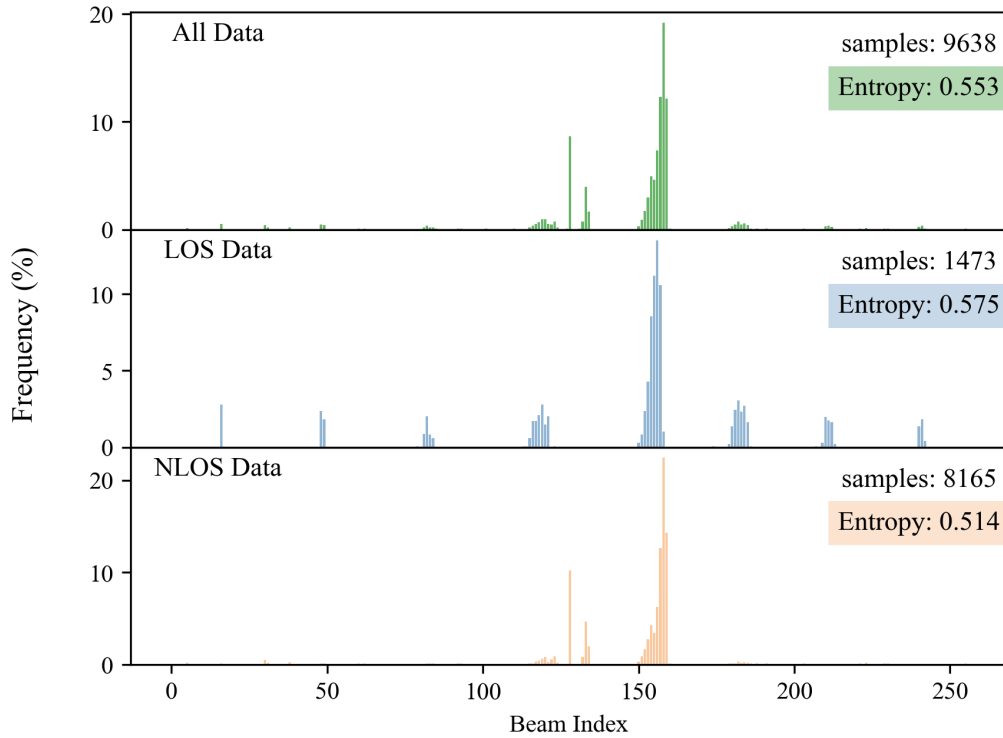


Figura 31: Distribuição de classes dos índices de feixe para os conjuntos de dados s009.

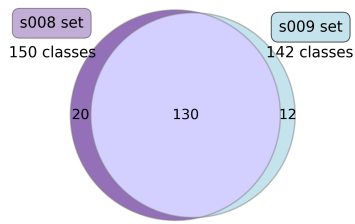
- $0,05 \leq |\Delta H| < 0,10$: Divergência moderada, afetando a confiança do modelo;
- $|\Delta H| \geq 0,10$: Mudança de domínio significativa que impacta a generalização [123, 124].

Para o cenário s008, H_{norm} é 0,625. A segmentação por tipo de propagação revela que o subconjunto LoS é mais determinístico ($H_{\text{norm}} = 0,556$) do que o subconjunto NLoS ($H_{\text{norm}} = 0,685$), como esperado devido aos efeitos de múltiplos caminhos. Curiosamente, s009 exibe uma entropia geral menor (0,553), com a condição NLoS apresentando a menor incerteza entre todos os casos ($H_{\text{norm}} = 0,514$). Esse resultado contra-intuitivo sugere que, em s009, o ambiente é limitado por geometrias semelhantes a corredores ou caminhos de reflexão dominantes que restringem a diversidade do feixe.

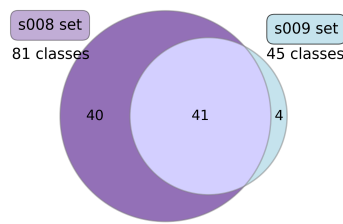
Além da análise de entropia, foi conduzido um estudo de sobreposição de classes (class overlap). Na figura 32 são identificadas as classes compartilhadas entre os datasets; índices de feixe exclusivos; compatibilidade entre espaços de rótulos.

Sob a condição de propagação ALL (Figura 32-a), observa-se um alto grau de sobreposição, com 130 classes comuns a ambos os conjuntos de dados. Essa interseção significativa, que representa 87% das classes de s008 e 92% das de s009, sugere uma forte semelhança estrutural nos padrões gerais de uso de feixes. Do ponto de vista da modelagem, esse espaço de rótulos compartilhado sustenta a viabilidade de estratégias de aprendizado por transferência ou treinamento conjunto. No entanto, a presença de classes exclusivas (20 em s008 e 12 em s009) evidencia caminhos de propagação específicos de cada cenário, ressaltando a necessidade de modelos capazes de lidar com rótulos não vistos no domínio de destino — uma característica crítica para a inferência em tempo real.

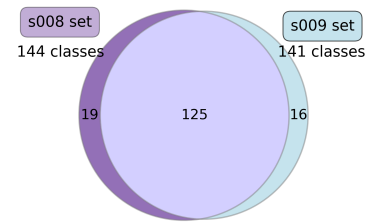
O cenário muda drasticamente sob a condição LoS (Figura 32-b). O espaço compartilhado reduz-se a apenas 41 classes, indicando que os padrões de propagação com visada direta (LoS)



(a) Todas as condições de propagação



(b) condições de propagação LoS



(c) condições de propagação NLoS

Figura 32: Diagrama de Venn mostrando a sobreposição das classes de índice de feixe nos conjuntos de dados s008 e s009 sob condição de propagação.

são bastante distintos entre os dois cenários. A assimetria é notável: s008 contém 40 classes LoS exclusivas, enquanto s009 possui apenas 4, provavelmente devido a diferenças na densidade e na distribuição geométrica dos receptores com visada direta. Essa sobreposição reduzida impõe um desafio significativo à generalização, sugerindo que um modelo treinado exclusivamente com dados LoS de um ambiente teria dificuldade de desempenho em outro, devido a uma incompatibilidade acentuada nos espaços de rótulos.

Por outro lado, a condição NLoS (Figura 32-c) mostra um retorno a um amplo espaço compartilhado, com 125 classes em comum. Essa alta sobreposição, comparável à da condição ALL, sugere que a propagação NLOS, impulsionada por reflexões e espalhamento, tende a gerar distribuições de feixes mais semelhantes e diversas em diferentes geometrias urbanas. Para o aprendizado de máquina, isso cria uma condição mais favorável para a generalização entre conjuntos de dados. O rico espaço de feixes compartilhado em cenários NLoS pode permitir que os modelos aprendam características mais robustas, melhorando o desempenho de previsão quando implantados em ambientes novos, porém estruturalmente semelhantes.

Em suma, essa análise de sobreposição revela que, embora os conjuntos de dados compartilhem uma parcela substancial de seus índices de feixes no geral, essa compatibilidade depende fortemente das condições de propagação. Embora o número de classes compartilhadas seja um pré-requisito para a generalização, ele não captura as disparidades em suas frequências de ocorrência.

Essa análise permitiu observar que determinadas condições de propagação apresentam reduzida interseção de classes, o que limita diretamente a transferência de conhecimento entre domínios. Passamos, portanto, a quantificar a divergência nessas distribuições compartilhadas.

11.2.5 Análise de Divergência entre Domínios

Na etapa seguinte, a análise foi expandida por meio da utilização de métricas clássicas e avançadas de distância entre distribuições. Foram avaliadas: divergência KL, distancia Wasserstein, Optimal Transport Dataset Distance (OTDD), KNN e KNN-class e Cross-Learning Score (CLS).

A utilização dessas métricas permitiu investigar o problema de *domain mismatch* sob diferentes perspectivas: divergência probabilística, distância geométrica, transporte ótimo e desempenho empírico entre domínios.

- **Divergência de Kullback–Leibler (KL):** Em análise de datasets, a KL permite avaliar o quanto a distribuição estatística de um conjunto de dados se distancia de uma distribu-

ição de referência, sendo particularmente útil para investigar mudanças de domínio (*dataset shift*), diferenças entre conjuntos de treinamento e teste, variações de características estatísticas e possíveis impactos na capacidade de generalização de modelos preditivos.

Considerando duas distribuições de probabilidade discretas, $(P(x))$ e $(Q(x))$, a divergência KL de (Q) em relação a (P) é definida como:

$$D_{KL}(P||Q) = \sum_{x \in X} P(x) \log \frac{P(x)}{Q(x)} \quad (95)$$

onde $(P(x))$ representa a distribuição de referência (ou distribuição real dos dados) e $(Q(x))$ representa a distribuição aproximada ou comparativa. O termo $(\log \frac{P(x)}{Q(x)})$ mede a diferença relativa entre as probabilidades atribuídas por ambas as distribuições para cada evento (x) . Dessa forma, a divergência KL representa a quantidade média de informação perdida ao utilizar (Q) para representar (P) .

No contexto de análise exploratória e avaliação de datasets, valores baixos de divergência KL indicam que os conjuntos analisados apresentam distribuições estatísticas semelhantes, sugerindo maior compatibilidade entre os dados. Por outro lado, valores elevados indicam uma alteração significativa na distribuição dos dados, podendo representar diferenças na composição das classes, mudança no padrão das características ou deslocamento do domínio (*domain shift*). Esse comportamento é relevante em aplicações de aprendizado supervisionado, pois modelos treinados em uma distribuição (P) podem apresentar degradação de desempenho quando aplicados a dados cuja distribuição (Q) apresenta grande divergência.

Em problemas de classificação, a divergência KL pode ser aplicada para comparar distribuições de classes. Essa análise permite identificar se dois datasets possuem estruturas estatísticas equivalentes ou se existem diferenças intrínsecas que podem afetar a transferência de conhecimento entre cenários. Além disso, quando associada a outras métricas de similaridade e complexidade dos dados, a KL fornece uma medida quantitativa para caracterizar a variabilidade dos datasets e auxiliar na interpretação dos resultados obtidos pelos modelos de aprendizado de máquina.

- **Wasserstein Distance:** A Distância de Wasserstein, também conhecida como *Earth Mover's Distance* (EMD), é uma métrica baseada na teoria de transporte ótimo (*Optimal Transport*) utilizada para quantificar a diferença entre duas distribuições de probabilidade. Diferentemente de métricas que realizam comparações pontuais entre distribuições, a Wasserstein considera a geometria do espaço de características, determinando o menor custo necessário para transportar a massa probabilística de uma distribuição até outra. Essa propriedade permite comparar distribuições mesmo quando apresentam pouca ou nenhuma sobreposição, tornando a métrica particularmente adequada para análise de mudanças de distribuição (*dataset shift*). Por se tratar de uma medida simétrica, a Distância de Wasserstein fornece um único valor de divergência para cada par de datasets, caracterizando apenas a discrepância estrutural entre suas distribuições. Consequentemente, a métrica não avalia explicitamente a capacidade de transferência ou a generalização cruzada entre os domínios de origem e destino

A formulação geral da Distância de Wasserstein entre duas distribuições de probabilidade

(P) e (Q) é definida como:

$$W_p(P,Q) = \left(\inf_{\gamma \in \Pi(P,Q)} \int_{X \times X} d(x,y)^p d\gamma(x,y) \right)^{\frac{1}{p}} \quad (96)$$

onde (p) representa a ordem da distância de Wasserstein, $(d(x,y))$ é a distância entre dois pontos (x) e (y) no espaço de características, e $(\Pi(P,Q))$ representa o conjunto de todas as possíveis distribuições de transporte (*couplings*) que transformam a distribuição (P) em (Q). A variável (γ) descreve como a massa probabilística deve ser redistribuída entre as duas distribuições, enquanto o termo integral representa o custo total desse transporte. A distância de Wasserstein busca, portanto, o menor custo possível para realizar essa transformação.

A interpretação intuitiva da Wasserstein é semelhante ao movimento de terra entre dois montes de areia: uma distribuição representa a configuração inicial da areia e a outra a configuração final. A distância mede a quantidade mínima de trabalho necessária para mover a massa da primeira distribuição até a segunda, considerando tanto a quantidade de massa deslocada quanto a distância percorrida. Dessa forma, distribuições próximas no espaço de características apresentam baixos valores de Wasserstein, enquanto distribuições significativamente diferentes resultam em valores elevados.

Assim, a Distância de Wasserstein constitui uma métrica poderosa para caracterização estatística de datasets, pois combina informações probabilísticas e geométricas, permitindo uma avaliação mais completa das diferenças entre distribuições de dados.

- **Optimal Transport Dataset Distance (OTDD):** A Optimal Transport Dataset Distance (OTDD) é uma métrica baseada na teoria de Transporte Ótimo (Optimal Transport – OT) desenvolvida para medir a distância entre dois conjuntos de dados considerando conjuntamente a distribuição das características (features) e a estrutura das classes associadas às amostras. De maneira análoga à Distância de Wasserstein, a OTDD é uma medida simétrica baseada em transporte ótimo e fornece um único valor para quantificar a divergência entre dois datasets. Embora incorpore informações tanto das características quanto dos rótulos, a métrica não modela explicitamente a direção da transferência entre domínios nem avalia a capacidade de generalização cruzada entre tarefas.

A OTDD utiliza a Distância de Wasserstein como o mecanismo fundamental para quantificar a discrepância entre distribuições. Enquanto a Distância de Wasserstein calcula o custo mínimo necessário para transportar massa probabilística entre duas distribuições considerando apenas a geometria do espaço de características, a OTDD estende esse conceito ao incorporar informações supervisionadas referentes aos rótulos das classes. Em outras palavras, a OTDD redefine a função de custo do problema de transporte ótimo de modo que o custo de transporte entre duas amostras dependa não apenas de sua proximidade no espaço de características, mas também da similaridade entre seus rótulos. Essa extensão permite capturar simultaneamente diferenças na distribuição dos atributos e na organização das classes, fornecendo uma caracterização mais abrangente da similaridade ou divergência entre datasets.

Na formulação da OTDD o objetivo é encontrar o menor custo necessário para transformar uma distribuição de dados em outra. Considerando dois datasets:

$$D_1 = (x_i, y_i) i = 1^n, \quad D_2 = (x_j, y_j) j = 1^m \quad (97)$$

onde (x) representa as características dos dados e (y) representa seus respectivos rótulos ou classes, a OTDD busca uma correspondência ótima entre as amostras dos dois datasets por meio de uma função de custo definida no espaço conjunto de características e classes. A distância OTDD pode ser representada de forma geral como:

$$OTDD(D_1, D_2) = \min_{\gamma \in \Pi(D_1, D_2)} \sum_{i,j} \gamma_{ij} C((x_i, y_i), (x_j, y_j)) \quad (98)$$

onde (γ) representa o plano ótimo de transporte entre os datasets, ($\Pi(D_1, D_2)$) representa o conjunto de possíveis distribuições de transporte entre as amostras dos dois conjuntos e ($C(\cdot, \cdot)$) corresponde à função de custo que mede a diferença entre pares de amostras considerando tanto a distância no espaço de características quanto a diferença entre seus rótulos. Uma característica fundamental da OTDD é a inclusão da informação de classe no cálculo da distância. Enquanto métricas como a Distância de Wasserstein tradicional avaliam principalmente a discrepância entre distribuições de características, a OTDD considera também a organização das classes dentro do espaço de dados. Dessa forma, dois datasets podem apresentar distribuições de atributos semelhantes, mas uma distribuição de classes diferente, resultando em uma maior distância OTDD. Essa propriedade torna a técnica especialmente relevante para problemas de aprendizado supervisionado. Em análise de datasets, a OTDD é utilizada para quantificar diferenças entre conjuntos de treinamento e teste, avaliar mudanças de domínio (domain shift), investigar transferência de aprendizado e identificar cenários onde um modelo pode apresentar redução de desempenho devido a diferenças estruturais nos dados. Valores baixos de OTDD indicam que os datasets possuem estruturas semelhantes em termos de características e classes, enquanto valores elevados indicam maior divergência entre os domínios analisados.

- **Métricas Baseadas em Vizinhança:** A métrica baseada no algoritmo k-Nearest Neighbors (KNN) com ($k=1$) é uma abordagem utilizada para avaliar a similaridade entre datasets por meio da proximidade entre suas amostras no espaço de características. Diferentemente de métricas baseadas em distribuições de probabilidade, como a divergência de Kullback-Leibler (KL) e a Distância de Wasserstein, essa técnica fundamenta-se na estrutura local dos dados, analisando a relação entre cada amostra e seu vizinho mais próximo. Seu princípio é que conjuntos de dados semelhantes tendem a apresentar elevada proximidade entre suas observações, enquanto datasets distintos exibem menor sobreposição e maior separação espacial.

A principal vantagem da abordagem baseada em KNN é que ela não requer a estimação explícita das distribuições de probabilidade, sendo aplicável a dados de alta dimensionalidade e distribuições complexas. Além disso, sua implementação é simples e apresenta baixo custo computacional para valores reduzidos de (k). Entretanto, por considerar apenas relações locais entre as amostras, a métrica pode ser sensível à densidade dos dados, à presença de ruído e à escolha da métrica de distância. Assim, a abordagem KNN com ($k=1$) é particularmente útil para detectar diferenças estruturais entre datasets, sendo amplamente empregada em tarefas de detecção de dataset shift, validação de domínio, comparação de distribuições empíricas e avaliação da similaridade entre conjuntos de treinamento e teste em aplicações de aprendizado de máquina.

- **Cross-Learning Score (CLS):** é uma métrica proposta para quantificar a similaridade entre conjuntos de dados supervisionados por meio da capacidade de generalização cruzada entre seus modelos preditivos. Diferentemente das métricas tradicionais, como a divergência de Kullback-Leibler (KL) e a Distância de Wasserstein, que avaliam predominantemente a distribuição das características, o CLS considera explicitamente a relação entre atributos e rótulos. A premissa central é que dois datasets são semelhantes quando um modelo treinado em um deles apresenta desempenho comparável ao modelo ótimo do próprio domínio ao ser aplicado ao outro conjunto de dados. Dessa forma, a métrica avalia diretamente a compatibilidade entre as estruturas de decisão dos datasets, tornando-se especialmente adequada para problemas de aprendizado supervisionado e aprendizado por transferência.

O cálculo do CLS baseia-se na diferença entre o desempenho dos modelos em seus respectivos domínios de origem e o desempenho obtido quando esses modelos são aplicados ao domínio oposto. Essa avaliação é realizada de forma bidirecional, considerando tanto a transferência do conjunto de origem para o conjunto alvo quanto o processo inverso. O valor final do CLS corresponde à média das discrepâncias de generalização observadas nas duas direções, produzindo uma medida simétrica da similaridade entre os datasets.

A Tabela 9 resume as métricas de divergência geométrica e de transferibilidade entre os cenários $s008$ e $s009$.

Métrica	ALL		LOS		NLOS	
	$\mathcal{D}_S \rightarrow \mathcal{D}_T$	$\mathcal{D}_T \rightarrow \mathcal{D}_S$	$\mathcal{D}_S \rightarrow \mathcal{D}_T$	$\mathcal{D}_T \rightarrow \mathcal{D}_S$	$\mathcal{D}_S \rightarrow \mathcal{D}_T$	$\mathcal{D}_T \rightarrow \mathcal{D}_S$
Divergência KL	0,190	0,191	0,406	0,866	0,483	0,560
KNN	4,00	6,40	6,78	24,40	6,93	4,90
KNN-class	7,68	5,47	31,47	7,28	10,34	6,63
Wasserstein	5,23		10,15		9,29	
OTDD	15,97		27,36		18,79	
[<i>feature</i>]						
OTDD	28,22		42,56		33,51	
[<i>feature + labels</i>]						
CLS	60,25%	50,66%	68,90%	28,60%	51,47%	20%
J. Manjarres	58,82%	48,01%	36,1%	28,75%	40,17%	24,87%
CLS						
J. Ruseckas	42,6%	38,1%	39,6%	7,39%	25,05%	21,56%
CLS						
B. Saleh						

$\mathcal{D}_S = s008, \mathcal{D}_T = s009$

Tabela 9: Métricas geométricas e de transferibilidade entre os domínios $s008$ e $s009$.

O Colapso Geométrico em LoS: Uma descoberta contraintuitiva surge ao analisar os ambientes de Visada Direta (LoS). Fisicamente, os canais LoS são determinísticos e proporcionam a maior precisão intradomínio (até 68,90% ao treinar e testar dentro do cenário $s008$). No entanto, como mostra a Tabela 9, o subconjunto LoS apresenta a maior divergência geométrica simétrica (OTDD = 42,57) e uma penalidade de distância 1-NN direcional elevada ($\mathcal{D}_S = s009 \rightarrow \mathcal{D}_T = s008 = 31,47$).

Isso revela um desalinhamento severo das características espaciais. Embora uma trajetória direta e desobstruída seja fácil de prever localmente, a geometria física (coordenadas GPS e nuvens de pontos LiDAR) dos caminhos diretos no traçado urbano de $s009$ não apresenta praticamente nenhuma semelhança topológica com o traçado viário de $s008$. Consequentemente, as fronteiras ótimas aprendidas em $s009$ não correspondem às coordenadas espaciais de $s008$, o que explica o colapso sistêmico na precisão LoS entre domínios (estagnada em $\approx 28\%$), apesar de sua simplicidade intradomínio.

O Paradoxo de Generalização NLoS:

Os subconjuntos isolados de Não-Visada-Direta (NLoS) apresentam um profundo paradoxo de generalização. Quando avaliados puramente com base em métricas do espaço de rótulos (Tabela 9), os domínios NLoS parecem estruturalmente compatíveis: eles compartilham a vasta maioria de seu suporte ativo ($|\mathcal{I}_{overlap}| = 125$ classes) e registram uma divergência KL relativamente baixa ($D_{KL} = 0,560$). Com base nessas métricas de rótulos e no enorme volume de treinamento disponível (8.165 amostras NLoS em $s009$), a transferência de domínio ($\mathcal{D}_S = s009 \rightarrow \mathcal{D}_T = s008$) deveria, teoricamente, ser bem-sucedida.

No entanto, a avaliação empírica revela um colapso catastrófico na acurácia, caindo para 20,00% para o framework WiSARD e 28,87% para J. Ruseckas. A análise voltada para diagnósticos geométricos (Tabela 9) resolve esse paradoxo. A pontuação OTDD simétrica dispara para 33,52, demonstrando que, embora os índices de feixe necessários permaneçam frequentes, as características multimodais de LiDAR e GPS sofrem uma severa mudança de covariável *covariate shift*. O complexo espalhamento por multicaminhos e as reflexões geradas pelo denso cânion urbano em $s009$ distorcem as características geométricas, tornando obsoletos no domínio $s008$ os mapeamentos do domínio de origem.

Essa incompatibilidade semântica é definitivamente confirmada pelas métricas direcionais. Enquanto os dados agregados (ALL) mantêm uma pontuação de aprendizado cruzado *Cross-Learning Score* - CLS bidirecional moderada (CLS = 55,45%), o isolamento dos dados NLoS revela um colapso drástico na transferibilidade (CLS = 35,73%) e uma penalidade elevada na classe KNN (10,34). Um CLS reduzido garante matematicamente que preditores ideais formulados no domínio de origem falharão ao generalizar para a distribuição de destino. Consequentemente, a barreira de generalização entre cenários não é um artefato de volume de dados ou capacidade do modelo insuficientes, mas sim uma consequência sistêmica do desalinhamento simétrico de características (alto OTDD) e da incompatibilidade assimétrica de fronteiras (baixo CLS e alto erro de classe KNN). Para corroborar fisicamente esse desalinhamento de características espaciais, a Figura 33 mapeia a distribuição geográfica dos receptores (Rx) em relação às coordenadas espaciais para ambos os cenários. A evidência visual revela um severo confinamento geográfico sob condições de LoS no cenário $s009$. Enquanto os receptores em $s008_{LoS}$ abrangem uma ampla faixa espacial ($x \in [745.22, 765.91], y \in [431.39, 676.89]$), as conexões LoS em $s009$ estão estritamente concentradas em uma caixa delimitadora espacial muito menor ($x \in [748.77, 763.55], y \in [579.74, 510.14]$), atuando efetivamente como um "zoom" espacial localizado. Os destaques em vermelho na Figura 33 marcam os limites de coordenadas superiores e inferiores que estão completamente ausentes do conjunto de dados $s009_{LoS}$. Esse truncamento geográfico explica a severa deficiência de suporte para classes fora do vocabulário (apenas 45 classes ativas em $s009_{LoS}$ em comparação com 81 em $s008_{LoS}$) e justifica visualmente a enorme penalidade direcional de classe KNN (31,47) e o alto valor de OTDD (42,57). Os modelos treinados em $s009_{LoS}$ sofrem uma falha catastrófica ao prever feixes ideais para as regiões de coordenadas não representadas de $s008_{LoS}$.

Os resultados mostraram que métricas baseadas em transporte ótimo capturam melhor

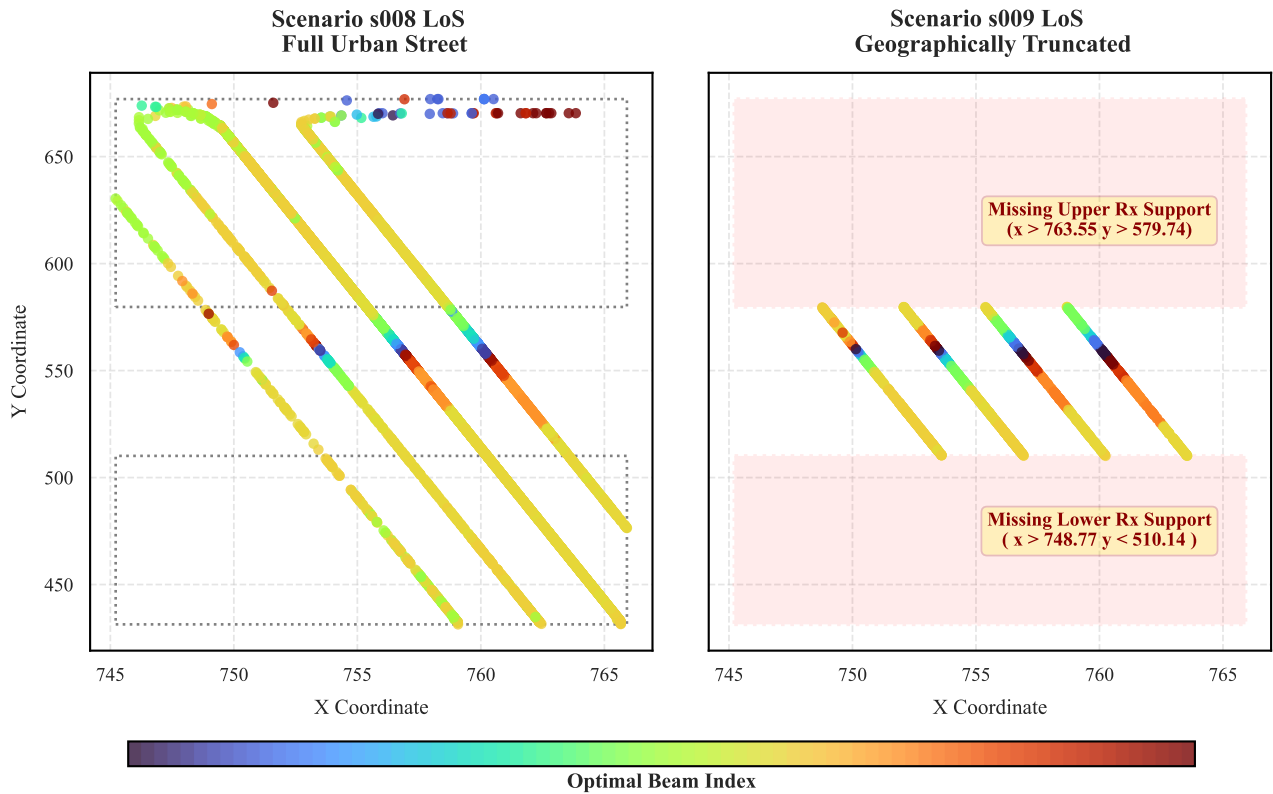


Figura 33: Distribuição geográfica dos receptores (Rx) nos cenários $s008$ e $s009$ sob propagação LoS. O cenário $s009_{LoS}$ (à direita) atua como um "zoom" espacial de $s008_{LoS}$ (à esquerda), resultando em uma grave deficiência de suporte. As caixas tracejadas em vermelho destacam as regiões superior ($y > 579$) e inferior ($y < 510$) não representadas em $s009_{LoS}$, o que explica o colapso catastrófico do suporte para classes fora do vocabulário ($|I_{overlap}| = 41$ de um total de 81 classes ativas).

diferenças estruturais globais; KL Divergence apresenta elevada sensibilidade local; CLS correlaciona diretamente divergência estatística e capacidade de generalização. Observou-se também que cenários aparentemente semelhantes em termos físicos podem apresentar distribuições estatísticas significativamente distintas, explicando parte da degradação observada nos modelos de ML.

11.3 Principais Resultados Obtidos

As atividades desenvolvidas até o momento permitiram identificar diversos aspectos importantes sobre o comportamento dos datasets utilizados na seleção de feixe para redes 6G:

- A análise anterior estabelece que o volume de dados, por si só, é um indicador pouco confiável da robustez do modelo na seleção de feixes em ondas milimétricas (mmWave) sob mudanças espaciais severas.
- A generalização dos modelos é fortemente influenciada pela estrutura estatística das distribuições de feixes;
- Cenários LoS e NLoS apresentam propriedades significativamente distintas;

- A multimodalidade dos dados (LiDAR + coordenadas) reduz parcialmente o impacto do desalinhamento entre domínios;
- Métricas baseadas em transporte ótimo mostraram maior robustez para caracterização de diferenças globais entre datasets;
- Métricas locais de vizinhança complementam a análise geométrica dos dados;
- O desempenho dos modelos depende não apenas da arquitetura de ML utilizada, mas também da compatibilidade estrutural entre os datasets.

Para sintetizar o poder preditivo da estrutura de diagnóstico proposta, avaliamos a consistência ordinal das métricas utilizando o coeficiente de correlação de postos de Kendall (τ). Conforme estabelecido na literatura sobre complexidade de dados, estatísticas baseadas em postos, como o τ de Kendall, são altamente adequadas para avaliar se métricas de complexidade conseguem prever ordinalmente a dificuldade de classificação entre diferentes domínios de conjuntos de dados.

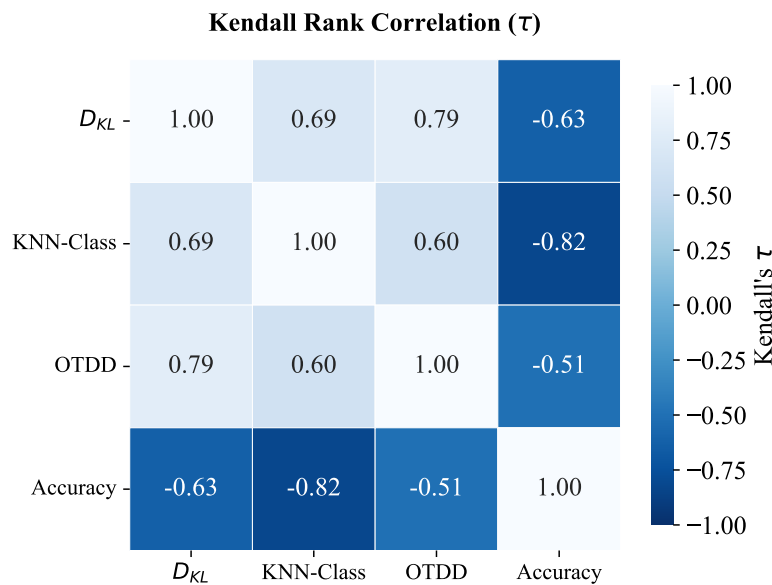


Figura 34: Mapa de calor da correlação de postos de Kendall (τ) detalhando as relações entre métricas de divergência geométricas/teórico-informacionais e a acurácia empírica fora da distribuição dos modelos avaliados.

A análise de correlação foi realizada usando as métricas KL, KNN-class, $OTDD_{[features+labels]}$ e acurácia no modelo WiSARD. Esta correlação valida matematicamente a capacidade preditiva das métricas geométricas e de teoria da informação em relação à degradação da acurácia empírica das arquiteturas avaliadas. A análise revela uma forte colinearidade positiva entre as métricas de divergência, com a divergência de Kullback-Leibler demonstrando alta correlação tanto com a OTDD ($\tau \approx 0,79$) quanto com a distância de classe KNN ($\tau \approx 0,69$). Isso confirma que mudanças marginais severas nos rótulos frequentemente ocorrem simultaneamente a distorções estruturais no espaço de características em ambientes de ondas milimétricas (mmWave).

Mais importante ainda, as métricas apresentam fortes correlações negativas com a acurácia de modelos de aprendizado de máquina em cenários *out-of-distribution* (fora da distribuição de treinamento). A distância de classe KNN direcional destaca-se como o preditor mais contundente de colapso de desempenho, registrando uma correlação negativa expressiva ($\tau = -0,82$) com a acurácia de classificação da arquitetura WiSARD.

As métricas de divergência global, a saber, KL ($\tau = -0,63$) e OTDD ($\tau = -0,51$), corroboram ainda mais essa tendência. Embora a OTDD capture efetivamente a mudança global de covariáveis gerada pelos cânions urbanos densos, a correlação extrema da penalidade KNN demonstra que é o entrelaçamento localizado das fronteiras multimodais que, em última análise, compromete fatalmente a capacidade de generalização do modelo.

Essa síntese estatística demonstra de forma definitiva que nossa estrutura centrada em dados atua como um indicador antecipado altamente confiável de falha na transferibilidade. Ao calcular *a priori* a penalidade de fronteira de classe KNN, a OTDD e a divergência KL, os orquestradores de redes 6G podem classificar com precisão a dificuldade esperada de transferir um modelo multimodal de seleção de feixes para um novo cânion urbano.

11.4 Conclusão e trabalhos futuros

As atividades conduzidas durante esta fase consolidaram uma importante linha de investigação centrada em dados para o problema de seleção de feixe em redes 6G. Foi desafiado o paradigma predominante centrado em modelos para a seleção de feixes (beam selection) em sistemas 6G de ondas milimétricas (mmWave), expondo as vulnerabilidades críticas de arquiteturas de aprendizado de máquina multimodal diante de mudanças no domínio espacial. Por meio de um protocolo rigoroso de avaliação bidirecional entre cenários, revelamos o "Paradoxo do Volume de Dados": evidências empíricas demonstraram que volumes massivos de amostras de treinamento, particularmente em ambientes densos de não visada direta (NLoS), não garantem a generalização para dados fora da distribuição (*out-of-distribution*), fazendo com que a precisão preditiva caia para níveis próximos ao acaso (*baseline* aleatório).

Para diagnosticar essa degradação catastrófica de desempenho sem depender de avaliações de modelo *post-hoc* computacionalmente custosas, propusemos uma estrutura abrangente de diagnóstico de dados que não requer treinamento. Nossa análise forense revelou que a generalização em cenários de visada direta (LoS) é limitada principalmente pelo truncamento geográfico, resultando em deficiências severas de suporte de vocabulário (*out-of-vocabulary*). Por outro lado, resolvemos o paradoxo NLoS ao demonstrar que, embora as distribuições marginais de rótulos possam parecer matematicamente compatíveis (apresentando baixa divergência de Kullback-Leibler e alta sobreposição de vocabulário ativo), as geometrias subjacentes das características multimodais sofrem mudanças severas de covariável (*covariate shifts*) induzidas por espalhamento complexo de múltiplos caminhos.

A integração da métrica de distância de conjuntos de dados via Transporte Ótimo (OTDD), da distância de classe 1-NN e do *Cross-Learning Score* (CLS) permitiu quantificar com sucesso esses desalinhamentos geométricos latentes. A robusta consistência ordinal de nossa estrutura, validada pela correlação de postos de Kendall ($\tau = -0,82$ para a penalidade espacial 1-NN), confirma matematicamente que a sobreposição de vizinhança local e o desalinhamento entre características e alvos entre domínios são os fatores determinantes que inviabilizam a transferibilidade do modelo, anulando efetivamente quaisquer vantagens decorrentes da profundidade da rede, da arquitetura ou do volume de amostras.

Em última análise, essa estrutura analítica centrada em dados serve como um importante

indicador precoce para a orquestração de redes 6G. Ao calcular essas métricas de divergência geométrica e de teoria da informação *a priori*, os operadores de sistemas podem estimar com precisão a dificuldade de transferência entre diferentes cânions urbanos, evitando os custos computacionais e energéticos proibitivos associados ao treinamento de redes neurais fadadas ao fracasso.

Trabalhos futuros focarão na integração dessas métricas diretamente em planos de controle nativos de IA para viabilizar a curadoria dinâmica de conjuntos de dados. Ao desenvolver um framework unificado de caracterização de datasets capaz de integrar múltiplas métricas em uma estrutura analítica única em tempo real. Os orquestradores de rede poderiam atuar como uma ferramenta preditiva para:

- estimar previamente dificuldades de generalização entre cenários;
- auxiliar na seleção de estratégias de treinamento e adaptação;
- orientar processos de coleta e balanceamento de dados;
- identificar cenários críticos para técnicas de aprendizado online;
- fornecer suporte para futuras estratégias de transferência de conhecimento entre domínios.

12 Gerenciamento de Recursos de Rádio em Redes MIMO Distribuídas Centradas no Usuário com TDD Dinâmico

Victor F. Monteiro, Yuri C. B. Silva, Igor M. Guerreiro
UFC

12.1 Introdução

O crescimento do tráfego em redes sem fio tem sido impulsionado por aplicações com requisitos heterogêneos e assimétricos, como *streaming* de vídeo e serviços em nuvem. Nesse cenário, diferentes equipamentos de usuário (*User Equipments* (UEs)) podem demandar simultaneamente elevadas taxas em enlace de descida (*Downlink* (DL)) e subida (*Uplink* (UL)), resultando em padrões de tráfego desbalanceados.

Sistemas convencionais utilizam duplexação por divisão no tempo (*Time Division Duplex* (TDD)), em que transmissões UL e DL seguem uma configuração fixa aplicada uniformemente a todos os UEs, reduzindo a eficiência espectral em cenários assimétricos [125]. Como alternativa, o *Dynamic TDD* (D-TDD) adapta dinamicamente a direção de transmissão conforme as condições de tráfego, aumentando a flexibilidade na alocação de recursos [126, 127, 128].

Além das limitações temporais, arquiteturas celulares convencionais também apresentam restrições no domínio espacial, uma vez que os UEs são tipicamente associados a uma única estação rádio base com antenas co-localizadas. Isso limita a diversidade espacial degradando a qualidade de serviço, especialmente nas bordas das células [129]. Nesse contexto, arquiteturas *Distributed Multi-input Multi-output* (D-MIMO), frequentemente associadas ao conceito de *cell-free Massive Multiple-Input Multiple-Output* (mMIMO), permitem que múltiplos pontos de acesso (*Access Points* (APs)) espalhados cooperem no atendimento aos UEs [130]. Para reduzir a complexidade da cooperação total, arquiteturas centradas no usuário restringem o atendimento de cada UE a subconjuntos de APs [131].

A integração entre D-TDD e arquiteturas centradas no usuário com D-MIMO amplia simultaneamente a flexibilidade temporal e espacial das redes sem fio, mas introduz desafios de gerenciamento de recursos de rádio (*Radio Resource Management* (RRM)), incluindo associação dinâmica entre UEs e APs, controle de potência, escalonamento de recursos e mitigação de interferências entre enlaces em direções opostas [125, 132].

Neste trabalho, considera-se a otimização conjunta da associação entre UEs e APs, da alocação de potência e do escalonamento de recursos em uma rede centrada no usuário com D-MIMO e D-TDD. O problema é inicialmente formulado como um problema de otimização não linear inteiro misto (*Mixed-Integer Nonlinear Problem* (MINLP)) e posteriormente reformulado como convexo inteiro misto (*Mixed-Integer Convex Optimization Problem* (MICP)). Além disso, avaliam-se os impactos da flexibilidade direcional do D-TDD e da cooperação espacial do D-MIMO em métricas como probabilidade de *outage* e taxa.

12.2 Modelo do Sistema

Considera-se um sistema composto por um conjunto $\mathcal{A}$ de A APs distribuídos, os quais atendem de forma cooperativa um conjunto $\mathcal{K}$ de K UEs com antena única. Os APs são conectados a uma unidade central de processamento (*Central Processing Unit* (CPU)), responsável pelo gerenciamento da rede, e possuem arranjos com N elementos de antena.

Um *slot* como unidade mínima de alocação. Cada *slot* possui duração t_{TTI} e é composto por um bloco de recurso (*Resource Block* (RB)) com largura de banda W_{RB} . A largura de banda total disponível é dada por $W = BW_{\text{RB}}$, correspondente ao conjunto $\mathcal{B}$ de B RBs.

Há T slots agrupados no conjunto $\mathcal{T}$. O UE k permanece em DL e UL durante T_k^{DL} e T_k^{UL} slots, respectivamente, tal que $T_k^{\text{DL}} + T_k^{\text{UL}} = T$. Além disso, definem-se os conjuntos $\mathcal{T}_k^{\text{DL}}$ e $\mathcal{T}_k^{\text{UL}}$, correspondentes aos intervalos nos quais o UE k opera em DL e UL, respectivamente.

Em um dado *slot* t , os UEs podem ser particionados nos conjuntos $\mathcal{K}_t^{\text{DL}}$ e $\mathcal{K}_t^{\text{UL}}$, que representam os UEs operando em DL e UL, respectivamente. De forma análoga, $\mathcal{A}_t^{\text{DL}}$ e $\mathcal{A}_t^{\text{UL}}$ representam os conjuntos de APs operando em DL e UL. Por fim, definem-se $\mathcal{A}_{k,t}$ e $\mathcal{K}_{a,t}$ no *slot* t como o conjunto de APs que atendem o UE k e o conjunto de UEs atendidos pelo AP a .

O canal $\mathbf{g}_{a,k,b,t}$ entre o AP a e o UE k no RB b e *slot* t é modelado como

$$\mathbf{g}_{a,k,b,t} = \sqrt{\beta_{a,k}} \mathbf{s}_{a,k,b,t} \in \mathbb{C}^N, \quad \text{com} \quad \beta_{a,k} = \left(\frac{d_0}{\max(d_{a,k}, d_{\min})} \right)^\alpha, \quad (99)$$

onde $\beta_{a,k}$ e $\mathbf{s}_{a,k,b,t} \sim \mathcal{CN}(\mathbf{0}, \mathbf{I}_N)$ são os desvanecimentos em larga e pequena escala, respectivamente, $d_{a,k}$ é a distância entre o AP a e o UE k , d_0 é uma distância de referência, $d_{\min}$ evita singularidades para pequenas distâncias e α é o expoente de perda de percurso.

De forma análoga, definem-se $h_{k,i,b,t} \in \mathbb{C}$ e $\mathbf{H}_{a',a,b,t} \in \mathbb{C}^{N \times N}$ como o canal entre os UEs k e i no *slot* t e no RB b , e o canal inter-AP entre os APs a' e a , respectivamente.

Com relação à transmissão e recepção de dados, no *slot* t e no RB b , o sinal $y_{k,b,t}^{\text{DL}} \in \mathbb{C}$ recebido pelo UE $k \in \mathcal{K}_t^{\text{DL}}$ em DL é dado por

$$y_{k,b,t}^{\text{DL}} = \underbrace{\sum_{a \in \mathcal{A}_{k,t}} \sqrt{\rho_{a,k,b,t}^{\text{DL}}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,k,b,t} s_{k,b,t}^{\text{DL}}}_{\text{Sinal desejado}} + \underbrace{\sum_{\substack{j \in \mathcal{K}_t^{\text{DL}} \\ j \neq k}} \sum_{a \in \mathcal{A}_{j,t}} \sqrt{\rho_{a,j,b,t}^{\text{DL}}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,j,b,t} s_{j,b,t}^{\text{DL}}}_{\text{Interf. multiusuário em DL}} + \underbrace{\sum_{i \in \mathcal{K}_t^{\text{UL}}} \sqrt{\rho_{i,b,t}^{\text{UL}}} h_{k,i,b,t} s_{i,b,t}^{\text{UL}}}_{\text{Interf. de usuários em UL}} + \underbrace{w_{k,b,t}}_{\text{Ruído}}. \quad (100)$$

O primeiro termo em (100) representa o sinal desejado transmitido pelos APs de $\mathcal{A}_{k,t} \subseteq \mathcal{A}_t^{\text{DL}}$. O escalar $\rho_{a,k,b,t}^{\text{DL}}$ é a potência de transmissão alocada pelo AP a ao UE k no *slot* t e no RB b , enquanto $\mathbf{v}_{a,k,b,t}$ é o vetor de *beamforming*. Além disso, $s_{k,b,t}^{\text{DL}} \sim \mathcal{CN}(0,1)$ é o símbolo transmitido ao UE k . O segundo termo é a interferência causada pelas transmissões aos demais UEs em DL. Já o terceiro termo é a interferência gerada pelos UEs em UL, onde $\rho_{i,b,t}^{\text{UL}}$ é a potência de transmissão do UE i , $h_{k,i,b,t}$ é o canal entre os UEs k e i , e $s_{i,b,t}^{\text{UL}} \sim \mathcal{CN}(0,1)$ é o símbolo transmitido em UL. Por fim, $w_{k,b,t} \sim \mathcal{CN}(0, \sigma^2)$ é o ruído gaussiano branco aditivo (*Additive White Gaussian Noise* (AWGN)) no receptor.

De forma análoga, o sinal $\mathbf{y}_{a,b,t}^{\text{UL}} \in \mathbb{C}^N$ recebido pelo AP $a \in \mathcal{A}_t^{\text{UL}}$ no *slot* t e no RB b é

$$\mathbf{y}_{a,b,t}^{\text{UL}} = \underbrace{\sum_{k \in \mathcal{K}_t^{\text{UL}}} \sqrt{\rho_{k,b,t}^{\text{UL}}} \mathbf{g}_{a,k,b,t} s_{k,b,t}^{\text{UL}}}_{\text{Sinais em UL}} + \underbrace{\sum_{k' \in \mathcal{K}_t^{\text{DL}}} \sum_{a' \in \mathcal{A}_{k',t}} \sqrt{\rho_{a',k',b,t}^{\text{DL}}} \mathbf{H}_{a',a,b,t} \mathbf{v}_{a',k',b,t} s_{k',b,t}^{\text{DL}}}_{\text{Interferência de transmissões em DL}} + \underbrace{\mathbf{w}_{a,b,t}}_{\text{Ruído}}, \quad (101)$$

onde $\mathbf{w}_{a,b,t} \sim \mathcal{CN}(\mathbf{0}, \sigma^2 \mathbf{1}_N)$ é o AWGN no AP a .

Cada AP $a \in \mathcal{A}_t^{\text{UL}}$ estima o símbolo transmitido pelo UE $k \in \mathcal{K}_{a,t}$ no RB b por meio do vetor linear de combinação $\mathbf{z}_{a,k,b,t}$, resultando em

$$\hat{s}_{a,k,b,t}^{\text{UL}} = \mathbf{z}_{a,k,b,t}^H \mathbf{y}_{a,b,t}^{\text{UL}}. \quad (102)$$

Consideramos *conjugate beamforming* na transmissão e recepção dos APs, isto é, $\mathbf{z}_{a,k,b,t} = \mathbf{v}_{a,k,b,t} = \hat{\mathbf{g}}_{a,k,b,t}$, sendo $\hat{\mathbf{g}}_{a,k,b,t}$ uma estimativa de $\mathbf{g}_{a,k,b,t}$.

As estimativas parciais obtidas pelos APs são posteriormente agregadas na CPU para decodificação do símbolo transmitido pelo UE k no RB b . O sinal decodificado é então

$$\hat{s}_{k,b,t}^{UL} = \sum_{a \in \mathcal{A}_{k,t}} \hat{s}_{a,k,b,t}^{UL}. \quad (103)$$

A relação sinal-interferência-ruído (*Signal to Interference plus Noise Ratio* (SINR)) de um UE $k \in \mathcal{K}_t^{UL}$ no RB b , considerando recepção conjunta pelos APs em $\mathcal{A}_{k,t} \subseteq \mathcal{A}_t^{UL}$, é dada por

$$\gamma_{k,b,t}^{UL} = \frac{\rho_{k,b,t}^{UL} \left| \sum_{a \in \mathcal{A}_{k,t}} \mathbf{z}_{a,k,b,t}^H \mathbf{g}_{a,k,b,t} \right|^2}{\underbrace{\sum_{\substack{j \in \mathcal{K}_t^{UL} \\ j \neq k}} \rho_{j,b,t}^{UL} \left| \sum_{a \in \mathcal{A}_{k,t}} \mathbf{z}_{a,k,b,t}^H \mathbf{g}_{a,j,b,t} \right|^2}_{\text{Interferência multiusuário em UL}} + \underbrace{\sum_{k' \in \mathcal{K}_t^{DL}} \left| \sum_{a \in \mathcal{A}_{k,t}} \sum_{a' \in \mathcal{A}_{k',t}} \sqrt{\rho_{a',k',b,t}^{DL}} \mathbf{z}_{a,k,b,t}^H \mathbf{H}_{a',a,b,t} \mathbf{v}_{a',k',b,t} \right|^2}_{\text{Interferência de APs em DL sobre UL}} + \underbrace{\sigma^2 \sum_{a \in \mathcal{A}_{k,t}} \|\mathbf{z}_{a,k,b,t}\|^2}_{\text{Ruído no receptor}}}. \quad (104)$$

De forma análoga, a SINR de um UE $k \in \mathcal{K}_t^{DL}$ no RB b é dada por

$$\gamma_{k,b,t}^{DL} = \frac{\left| \sum_{a \in \mathcal{A}_{k,t}} \sqrt{\rho_{a,k,b,t}^{DL}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,k,b,t} \right|^2}{\underbrace{\sum_{\substack{j \in \mathcal{K}_t^{DL} \\ j \neq k}} \left| \sum_{a \in \mathcal{A}_{j,t}} \sqrt{\rho_{a,j,b,t}^{DL}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,j,b,t} \right|^2}_{\text{Interferência multiusuário em DL}} + \underbrace{\sum_{i \in \mathcal{K}_t^{UL}} \rho_{i,b,t}^{UL} |h_{k,i,b,t}|^2}_{\text{Interferência de UEs em UL}} + \underbrace{\sigma^2}_{\text{Ruído térmico}}}. \quad (105)$$

Os numeradores das Eqs. (104) e (105) representam a combinação coerente dos sinais recebidos/transmitidos por múltiplos APs, respectivamente, enquanto os denominadores incorporam os efeitos de interferência multiusuário e cruzada entre transmissões em UL e DL.

Então, a quantidade de bits transmitidos ou recebidos pelo UE k no RB b e no *slot* t é

$$Q_{k,b,t} = \begin{cases} Q_{k,b,t}^{DL} = W_{\text{RB}} \log_2 (1 + \gamma_{k,b,t}^{DL}) t_{\text{TTI}}, & \text{para } t \in \mathcal{T}_k^{DL}, \\ Q_{k,b,t}^{UL} = W_{\text{RB}} \log_2 (1 + \gamma_{k,b,t}^{UL}) t_{\text{TTI}}, & \text{para } t \in \mathcal{T}_k^{UL}. \end{cases} \quad (106)$$

12.3 Formulação do Problema

O problema considerado pode ser formulado conforme (107). O objetivo, dado por (107a), consiste em maximizar a soma total de bits transmitidos e recebidos por todos os UEs ao longo dos *slots*, respeitando restrições de potência, requisitos mínimos de *Quality of Service* (QoS) e consistência de operação em TDD.

As Eqs. (107b) e (107c) impõem restrições de consistência TDD, impedindo que APs e UEs operem simultaneamente em DL e UL. Para isso, definem-se as variáveis binárias $y_{a,k,b,t}^{DL}, y_{a,k,b,t}^{UL} \in \{0,1\}$, que indicam se o AP a e o UE k estão conectados em DL ou UL, respectivamente, no *slot* t e no RB b . As Eqs. (107d) a (107g) definem as restrições de potência, em que $P_a^{\max}$ e $P_k^{\max}$ representam as potências máximas de transmissão do AP a e do UE k , respectivamente. A Eq. (107d) impõe que potência em DL só seja alocada quando AP e UE estão conectados, enquanto (107e) limita a potência total transmitida por cada AP. De forma análoga, (107f) garante que a potência em UL em determinado RB seja nula quando o recurso não está alocado ao UE, enquanto (107g) limita a potência total transmitida por cada UE. As Eqs. (107h) e (107i) representam as restrições de QoS, impondo quantidades mínimas de bits transmitidos e recebidos em UL e DL, respectivamente.

$$\max_{\rho_{a,k,b,t}^{DL}, \rho_{k,b,t}^{UL}, y_{a,k,b,t}^{DL}, y_{a,k,b,t}^{UL}} \sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}} Q_{k,b,t} \quad (107a)$$

$$\text{sujeito a } \left(\sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} y_{a,k,b,t}^{DL} \right) \left(\sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} y_{a,k,b,t}^{UL} \right) = 0, \quad \forall a \in \mathcal{A}, \forall t \in \mathcal{T}, \quad (107b)$$

$$\left(\sum_{a \in \mathcal{A}} \sum_{b \in \mathcal{B}} y_{a,k,b,t}^{DL} \right) \left(\sum_{a \in \mathcal{A}} \sum_{b \in \mathcal{B}} y_{a,k,b,t}^{UL} \right) = 0, \quad \forall k \in \mathcal{K}, \forall t \in \mathcal{T}, \quad (107c)$$

$$\rho_{a,k,b,t}^{DL} \leq y_{a,k,b,t}^{DL} P_a^{\max}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (107d)$$

$$\sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} \rho_{a,k,b,t}^{DL} \leq P_a^{\max}, \quad \forall a \in \mathcal{A}, \forall t \in \mathcal{T}, \quad (107e)$$

$$\rho_{k,b,t}^{UL} \leq \left(\sum_{a \in \mathcal{A}} y_{a,k,b,t}^{UL} \right) P_k^{\max}, \quad \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (107f)$$

$$\sum_{b \in \mathcal{B}} \rho_{k,b,t}^{UL} \leq P_k^{\max}, \quad \forall k \in \mathcal{K}, \forall t \in \mathcal{T}, \quad (107g)$$

$$\sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}_k^{DL}} Q_{k,b,t} \geq Q_k^{DL, \min}, \quad \forall k \in \mathcal{K}, \quad (107h)$$

$$\sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}_k^{UL}} Q_{k,b,t} \geq Q_k^{UL, \min}, \quad \forall k \in \mathcal{K}, \quad (107i)$$

$$\rho_{a,k,b,t}^{DL}, \rho_{k,b,t}^{UL} \geq 0, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (107j)$$

$$y_{a,k,b,t}^{DL}, y_{a,k,b,t}^{UL} \in \{0,1\}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}. \quad (107k)$$

12.4 Reformulação Proposta

Eq. (107) é um MINLP não convexo devido à presença de variáveis de otimização em: i) produtos bilineares em (107b) e (107c); ii) denominadores das expressões de SINR; e iii) acoplamentos não convexos entre potências de transmissão na combinação coerente dos sinais em DL no numerador de (105). Para obter um MICP, consideram-se: i) a linearização das restrições de consistência TDD; ii) a ortogonalização em frequência de UEs; e iii) *clusters* de APs com parametrização convexa de potência. O problema reformulado é apresentado em (111).

Linearização de Restrições de Consistência do TDD As Eqs. (111b)–(111f) linearizam as restrições de consistência TDD (Eqs. (107b) e (107c)) por meio das variáveis binárias auxiliares $s_{i,t}^{DL}, s_{i,t}^{UL} \in \{0,1\}$, em que $i \in \mathcal{A} \cup \mathcal{K}$. Essas variáveis indicam se o nó i opera em DL ou UL no *slot* t . Eq. (111b) garante que APs e UEs operem em apenas uma direção de transmissão por *slot*, enquanto (111c)–(111f) relacionam as variáveis de alocação e as direções de transmissão selecionadas. Garantindo operação na mesma direção entre APs e UEs conectados.

Ortogonalização de UEs na Frequência As expressões de SINR em (104) e (105) são não convexas porque os termos de interferência dependem das potências de transmissão de outros UEs. Uma abordagem clássica para regularizar o problema consiste em impor ortogonalização em frequência entre UEs, de forma que, em cada *slot*, cada RB possa ser utilizado por no máximo um UE e em apenas uma direção de transmissão (DL ou UL).

As Eqs. (111g)–(111i) impõem ortogonalização em frequência entre UEs por meio das variáveis binárias $x_{k,b,t}^{DL}, x_{k,b,t}^{UL} \in \{0,1\}$, que indicam se o RB b é alocado ao UE k em DL ou UL. A Eq. (111g) garante que cada RB seja utilizado por no máximo um UE em cada *slot*, enquanto (111h) e (111i) realizam o acoplamento entre alocação de RB e associação AP-UE. Essas duas garantem que mais de um AP podem servir simultaneamente um mesmo UE no

mesmo RB. Com isso, eliminam-se os termos de interferência inter-UE na SINR e Eqs. (104) e (105) tornam-se, respectivamente:

$$\gamma_{k,b,t}^{UL} = \frac{\rho_{k,b,t}^{UL} \left| \sum_{a \in \mathcal{A}_{k,t}} \mathbf{z}_{a,k,b,t}^H \mathbf{g}_{a,k,b,t} \right|^2}{\sigma^2 \sum_{a \in \mathcal{A}_{k,t}} \|\mathbf{z}_{a,k,b,t}\|^2} \quad \text{e} \quad \gamma_{k,b,t}^{DL} = \frac{\left| \sum_{a \in \mathcal{A}_{k,t}} \sqrt{\rho_{a,k,b,t}^{DL}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,k,b,t} \right|^2}{\sigma^2} \quad (108)$$

Clusterização de APs e Parametrização de Potência Para evitar o acoplamento não convexo entre potências de transmissão dos APs em (108), considera-se um conjunto de L *clusters* de APs candidatos $\{\mathcal{C}_{k,1}, \dots, \mathcal{C}_{k,L}\}$ para cada UE k , em que $\mathcal{C}_{k,\ell} \subseteq \mathcal{A}$. Define-se ainda $c_{a,k,\ell}$ como uma constante binária que indica se o AP a pertence ao *cluster* $\mathcal{C}_{k,\ell}$. Considera-se que cada *cluster* contém os ℓ APs com maiores coeficientes de desvanecimento em larga escala em relação ao UE k , embora outras estratégias possam ser utilizadas sem alterar a formulação proposta.

As variáveis binárias $u_{k,\ell,t} \in \{0,1\}$ indicam se o *cluster* $\mathcal{C}_{k,\ell}$ é selecionado para atender o UE k no *slot* t . A Eq. (111n) garante que no máximo um *cluster* seja selecionado por UE e *slot*.

A potência total transmitida ao UE k no RB b e *slot* t quando o *cluster* ℓ é selecionado é representada pelas variáveis contínuas $p_{k,b,t,\ell}^{DL} \geq 0$. As Eqs. (111o) e (111p) garantem que essa potência seja nula quando o *cluster* ou o RB correspondente não forem selecionados.

A distribuição da potência entre os APs do *cluster* é realizada por coeficientes fixos $\eta_{a,k,\ell}$, tais que $\eta_{a,k,\ell} \geq 0$, $\eta_{a,k,\ell} = 0$ para $a \notin \mathcal{C}_{k,\ell}$ e $\sum_{a \in \mathcal{C}_{k,\ell}} \eta_{a,k,\ell} = 1$. Nesta implementação, considera-se

$$\eta_{a,k,\ell} = \frac{\beta_{a,k}}{\sum_{a' \in \mathcal{C}_{k,\ell}} \beta_{a',k}}. \quad (109)$$

Com isso, a potência transmitida por cada AP pode ser escrita como uma função afim das variáveis contínuas de otimização, conforme (111q), enquanto (111r) garante que apenas APs pertencentes ao *cluster* selecionado participem da transmissão.

Eq. (111n) garante que só um *cluster* ℓ^* pode estar ativo por *slot*, logo, de (111o), existe no máximo uma variável $p_{k,b,t,\ell}^{DL} \neq 0$ ($p_{k,b,t,\ell}^{DL} = 0 \forall \ell \neq \ell^*$). Substituindo (111q) em (108), o somatório em ℓ se resume a um termo e este é independente de a , saindo do somatório em a . De forma equivalente, considerar apenas p_{k,b,t,ℓ^*}^{DL} ou somar em ℓ é igual, ou seja

$$\gamma_{k,b,t}^{DL} = \frac{\left| \sum_{a \in \mathcal{A}_{k,t}} \sqrt{\sum_{\ell=1}^L c_{a,k,\ell} \eta_{a,k,\ell} p_{k,b,t,\ell}^{DL}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,k,b,t} \right|^2}{\sigma^2} = \frac{p_{k,b,t,\ell^*}^{DL} \left| \sum_{a \in \mathcal{A}_{k,t}} \sqrt{c_{a,k,\ell^*} \eta_{a,k,\ell^*}} \mathbf{g}_{a,k,b,t}^H \mathbf{v}_{a,k,b,t} \right|^2}{\sigma^2} = \frac{\sum_{\ell=1}^L G_{k,b,t,\ell}^{DL} p_{k,b,t,\ell}^{DL}}{\sigma^2}. \quad (110)$$

Assim, a sequência de transformações proposta (linearização da consistência TDD, ortogonalização em frequência e parametrização convexa de potência baseada em clusterização) elimina todas as fontes de não convexidade presentes na formulação original. Assim, a SINR em DL torna-se afim nas variáveis contínuas de potência, enquanto as expressões de taxa tornam-se côncavas. Além disso, a detecção em UL permanece convexa sob alocação ortogonal de RBs,

eliminando os acoplamentos não lineares entre potências específicas dos APs. Portanto, a formulação resultante em (111) corresponde a um MICP.

$$\max_{\substack{y_{a,k,b,t}^{DL}, y_{a,k,b,t}^{UL}, \\ s_{i,t}^{DL}, s_{i,t}^{UL}, x_{k,b,t}^{DL}, x_{k,b,t}^{UL}, \\ u_{k,\ell,t}, p_{k,b,t,\ell}^{DL}, p_{k,b,t,\ell}^{UL}}} \sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}} Q_{k,b,t} \quad (111a)$$

sujeito a Consistência TDD nos UEs e APs:

$$s_{i,t}^{DL} + s_{i,t}^{UL} \leq 1, \quad \forall i \in \mathcal{A} \cup \mathcal{K}, \forall t \in \mathcal{T}, \quad (111b)$$

$$y_{a,k,b,t}^{DL} \leq s_{a,t}^{DL}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111c)$$

$$y_{a,k,b,t}^{UL} \leq s_{a,t}^{UL}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111d)$$

$$y_{a,k,b,t}^{DL} \leq s_{k,t}^{DL}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111e)$$

$$y_{a,k,b,t}^{UL} \leq s_{k,t}^{UL}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111f)$$

Ortogonalização em frequência:

$$\sum_{k \in \mathcal{K}} x_{k,b,t}^{DL} + \sum_{k \in \mathcal{K}} x_{k,b,t}^{UL} \leq 1, \quad \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111g)$$

$$y_{a,k,b,t}^{DL} \leq x_{k,b,t}^{DL}, \quad \forall k \in \mathcal{K}, \forall a \in \mathcal{A}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111h)$$

$$y_{a,k,b,t}^{UL} \leq x_{k,b,t}^{UL}, \quad \forall k \in \mathcal{K}, \forall a \in \mathcal{A}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111i)$$

Restrições de potência:

$$\rho_{a,k,b,t}^{DL} \leq y_{a,k,b,t}^{DL} \cdot P_a^{\max}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111j)$$

$$\sum_{k \in \mathcal{K}} \sum_{b \in \mathcal{B}} \rho_{a,k,b,t}^{DL} \leq P_a^{\max}, \quad \forall a \in \mathcal{A}, \forall t \in \mathcal{T}, \quad (111k)$$

$$\rho_{k,b,t}^{UL} \leq \left(\sum_{a \in \mathcal{A}} y_{a,k,b,t}^{UL} \right) P_k^{\max}, \quad \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111l)$$

$$\sum_{b \in \mathcal{B}} \rho_{k,b,t}^{UL} \leq P_k^{\max}, \quad \forall k \in \mathcal{K}, \forall t \in \mathcal{T}, \quad (111m)$$

Seleção de *cluster* e acoplamento de potência:

$$\sum_{\ell=1}^L u_{k,\ell,t} \leq 1, \quad \forall k \in \mathcal{K}, \forall t \in \mathcal{T}, \quad (111n)$$

$$p_{k,b,t,\ell}^{DL} \leq u_{k,\ell,t} A P_a^{\max}, \quad \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \forall \ell \in \{1, 2, \dots, L\}, \quad (111o)$$

$$p_{k,b,t,\ell}^{UL} \leq x_{k,b,t}^{UL} A P_a^{\max}, \quad \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \forall \ell \in \{1, 2, \dots, L\}, \quad (111p)$$

$$\rho_{a,k,b,t}^{DL} = \sum_{\ell=1}^L c_{a,k,\ell} \eta_{a,k,\ell} p_{k,b,t,\ell}^{DL}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111q)$$

$$y_{a,k,b,t}^{DL} \leq \sum_{l=1}^L c_{a,k,\ell} u_{k,\ell,t}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \quad (111r)$$

Restrições de QoS (bits em DL/UL):

$$\sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}_k^{DL}} Q_{k,b,t}^{DL} \geq Q_k^{DL, \min}, \quad \forall k \in \mathcal{K}, \quad (111s)$$

$$\sum_{b \in \mathcal{B}} \sum_{t \in \mathcal{T}_k^{UL}} Q_{k,b,t}^{UL} \geq Q_k^{UL, \min}, \quad \forall k \in \mathcal{K}, \quad (111t)$$

Domínio das variáveis:

$$y_{a,k,b,t}^{DL}, y_{a,k,b,t}^{UL}, s_{i,t}^{DL}, s_{i,t}^{UL}, x_{k,b,t}^{DL}, x_{k,b,t}^{UL}, u_{k,\ell,t} \in \{0, 1\}, \quad \forall a \in \mathcal{A}, \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \forall \ell \in \{1, 2, \dots, L\}, \quad (111u)$$

$$\rho_{k,b,t}^{UL}, p_{k,b,t,\ell}^{DL} \geq 0, \quad \forall k \in \mathcal{K}, \forall b \in \mathcal{B}, \forall t \in \mathcal{T}, \forall \ell \in \{1, 2, \dots, L\} \quad (111v)$$

12.5 Soluções de Referência

Nesta seção, são definidas soluções de referência para avaliar os ganhos proporcionados pelo D-TDD e pela operação centrada no usuário com D-MIMO. Elas são obtidas como casos

particulares de (111), por meio da desativação da cooperação multi-AP, do D-TDD, ou de ambos.

12.5.1 Remoção da Operação Centrada no Usuário com D-MIMO

Conectividade única em DL Em (111), a operação centrada no UE com D-MIMO no DL é representada por meio de *clusters* candidatos $\mathcal{C}_{k,\ell}$ associados a cada UE k . A conectividade única em DL é obtida restringindo cada *cluster* a conter exatamente um AP. Assim, as variáveis de seleção de *cluster* passam apenas a indicar qual AP atende cada UE, eliminando a cooperação multi-AP sem necessidade de restrições adicionais.

Conectividade única em UL No UL, a conectividade única não decorre naturalmente de (111), uma vez que um mesmo UE pode ser servido simultaneamente por múltiplos APs. Assim, introduzem-se variáveis binárias $z_{a,k,t}^{UL}$, em que $z_{a,k,t}^{UL} = 1$ indica que o UE k está associado ao AP a no *slot* t . A conectividade única é então imposta por

$$\sum_{a=1}^A z_{a,k,t}^{UL} \leq 1, \quad \forall k \in \mathcal{K}, \forall t \in \mathcal{T}, \quad \text{e} \quad y_{a,k,b,t}^{UL} \leq z_{a,k,t}^{UL}, \quad \forall a,k,b,t, \quad (112)$$

garantindo que um mesmo UE não possa ser associado a diferentes APs em distintos RBs em um mesmo *slot*.

12.5.2 Remoção do D-TDD (TDD Estático)

No D-TDD, cada AP e UE pode selecionar independentemente sua direção de transmissão a cada *slot*. Para modelar TDD estático, impõe-se uma direção de transmissão global para toda a rede em cada *slot*. Para isso, define-se uma variável binária z_t , em que $z_t = 1$ indica operação em DL e $z_t = 0$ operação em UL. As variáveis de direção em (111) são então restringidas por

$$s_{i,t}^{DL} \leq z_t, \quad s_{i,t}^{UL} \leq 1 - z_t, \quad \forall i \in \mathcal{A} \cup \mathcal{K}, \forall t \in \mathcal{T}, \quad (113)$$

forçando todos os APs e UEs a operarem na mesma direção em cada *slot*. As demais restrições da formulação em (111) permanecem inalteradas.

12.5.3 Configurações de Referência

A partir das especializações descritas anteriormente, definem-se três configurações de referência.

A primeira corresponde a um sistema com TDD estático e sem operação centrada no usuário com D-MIMO. Nesse caso, os *clusters* são restritos a conjuntos com único AP e a direção de transmissão é globalmente definida em cada *slot* conforme (113), caracterizando uma rede convencional com conectividade única.

A segunda configuração considera TDD estático com operação centrada no usuário com D-MIMO. Nesse caso, a cooperação multi-AP é mantida, porém todos os nós operam na mesma direção de transmissão em cada *slot*. Essa configuração permite isolar os ganhos proporcionados pela operação centrada no usuário com D-MIMO.

A terceira configuração corresponde ao D-TDD sem operação centrada no usuário com D-MIMO. Nesse caso, os UEs operam com conectividade única, enquanto APs e UEs podem selecionar independentemente suas direções de transmissão. Essa configuração permite avaliar os ganhos proporcionados exclusivamente pelo D-TDD.

Tabela 10: Parâmetros de simulação do sistema com D-MIMO e D-TDD.

Parâmetro	Valor
Número de APs (A)	4
Antenas por AP (N)	2
Área de simulação	200 m $\times$ 200 m
Número de <i>slots</i> (T)	5
Largura de banda do RB (W_{RB})	180 kHz
Duração do <i>slot</i> (t_{TTI})	1 ms
Potência do ruído (σ^2)	10^{-9} W
Expoente de perda de percurso (α)	3.7
Distância de referência (d_0)	1 m
Distância mínima ($d_{\min}$)	1 m
Potência de transmissão max. de APs e UEs ($P_a^{\max}$ e $P_k^{\max}$)	24 dBm
Realizações de Monte Carlo	100

12.6 Avaliação de Desempenho

Avaliam-se quatro arquiteturas: i) conectividade única + TDD estático; ii) operação centrada no usuário com D-MIMO + TDD estático; iii) conectividade única + D-TDD; e iv) operação centrada no usuário com D-MIMO + D-TDD. O objetivo é quantificar os impactos da cooperação espacial e da flexibilidade direcional.

A Tabela 10 apresenta os principais parâmetros de simulação. Os APs são posicionados em grade regular em uma área de 200 m $\times$ 200 m, enquanto os UEs são distribuídos uniformemente. Salvo indicação em contrário, variam-se K , B e os requisitos mínimos de QoS.

Na arquitetura centrada no usuário com D-MIMO, os *clusters* são definidos a partir dos maiores coeficientes $\beta_{a,k}$, considerando $\alpha = 3.7$ e $d_{\min} = d_0 = 1$ m. Especificamente, consideram-se 3 *clusters* por UE k : $C_{k,1} = \{a_{(1)}\}$, $C_{k,2} = \{a_{(1)}, a_{(2)}\}$, $C_{k,3} = \{a_{(1)}, a_{(2)}, a_{(3)}\}$, em que $a_{(i)}$ representa o AP com o i -ésimo maior coeficiente $\beta_{a,k}$.

Eq. (111) é resolvida com CVXPY e MOSEK. Para cada configuração, são realizadas 100 simulações de Monte Carlo. Considera-se *outage* quando o problema se torna inviável devido à violação das restrições de QoS. As estatísticas de taxa são calculadas apenas para realizações em que todas as arquiteturas são viáveis, reportando-se a taxa média e os percentis 10 e 90.

A Fig. 35 ilustra uma solução em cenário com operação centrada no UE com D-MIMO e D-TDD. Neste exemplo, consideram-se $K = 3$ UEs e $B = 2$ RBs. Cada subfigura apresenta as associações entre APs e UEs e a alocação de RBs em um dado *slot*. Linhas azuis contínuas representam transmissões em DL, enquanto linhas vermelhas tracejadas representam transmissões em UL. O rótulo “bx” indica o RB utilizado na transmissão, com $x \in \{0,1\}$.

Os três UEs possuem diferentes requisitos mínimos de taxa em DL e UL. Os requisitos mínimos em DL são 0.1, 0.5 e 0.7 Mbps, enquanto em UL são 0.5, 0.1 e 0.7 Mbps para os UEs 0, 1 e 2, respectivamente. Assim, cada UE apresenta um perfil de tráfego distinto entre as direções de transmissão.

A figura evidencia duas características principais da arquitetura proposta. Primeiramente, o *solver* determina dinamicamente a direção de transmissão de cada AP em cada *slot*. Por exemplo, em $t = 1$, o AP 3 opera em UL, enquanto os demais operam em DL; já em $t = 3$, os APs 0 e 1 operam em DL, enquanto os APs 2 e 3 operam em UL. Esses casos ilustram a

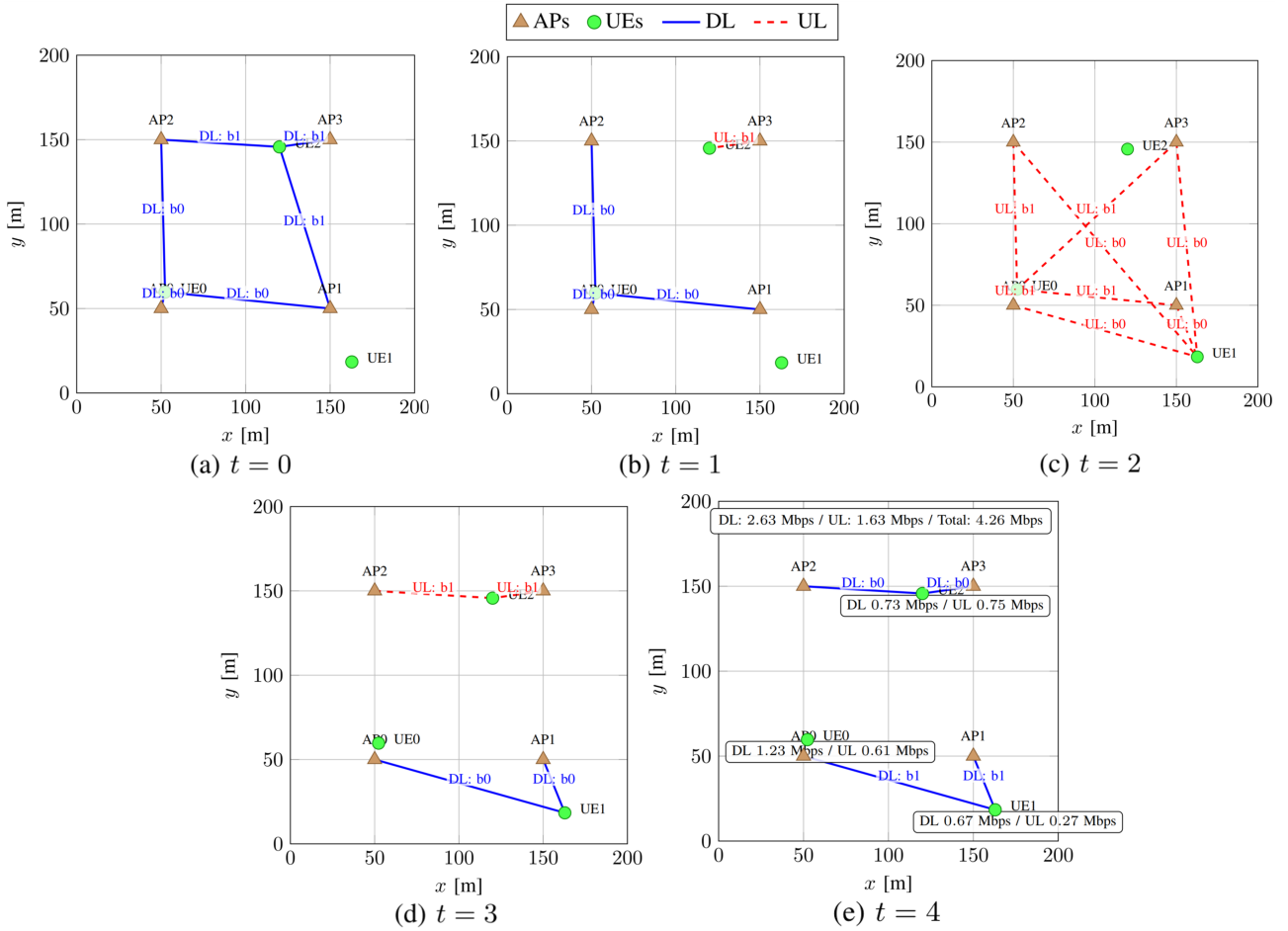


Figura 35: Associação entre APs e UEs e alocação de RBs ao longo do tempo em uma rede centrada no usuário com D-MIMO e D-TDD.

flexibilidade proporcionada pelo D-TDD, permitindo operação simultânea em direções opostas.

Além disso, a figura também evidencia a operação centrada no usuário com D-MIMO, em que múltiplos APs podem servir simultaneamente um mesmo UE utilizando o mesmo RB. Isso pode ser observado por enlaces com o mesmo índice de RB conectando diferentes APs ao mesmo UE, caracterizando transmissão cooperativa no mesmo recurso de frequência. Exemplo, em $t = 3$, o UE 1 é servido no DL no RB 0 (“DL:b0”) pelos APs 0 e 1.

A última subfigura apresenta ainda as taxas obtidas por cada UE em DL e UL, bem como as taxas agregadas em DL, UL e total do sistema. Todas as restrições mínimas de QoS são satisfeitas, enquanto os recursos remanescentes são alocados de forma oportunista para maximizar a vazão total da rede.

A. Impacto dos Requisitos de QoS Figs. 36a e 36d apresentam a probabilidade de *outage* e a taxa do sistema em função do requisito de QoS, considerando $B = 3$ RBs e $K = 4$ UEs. Observa-se aumento do *outage* com o crescimento dos requisitos de taxa, uma vez que maiores taxas demandam quantidades maiores de potência e/ou RBs. Como os recursos disponíveis são limitados, a região viável de (111) reduz-se à medida que o requisito de QoS aumenta.

A arquitetura centrada no usuário com D-MIMO e D-TDD apresenta consistentemente o menor *outage*, explorando simultaneamente diversidade espacial e flexibilidade direcional. Além disso, observa-se que conectividade única com D-TDD supera a operação centrada no usuário

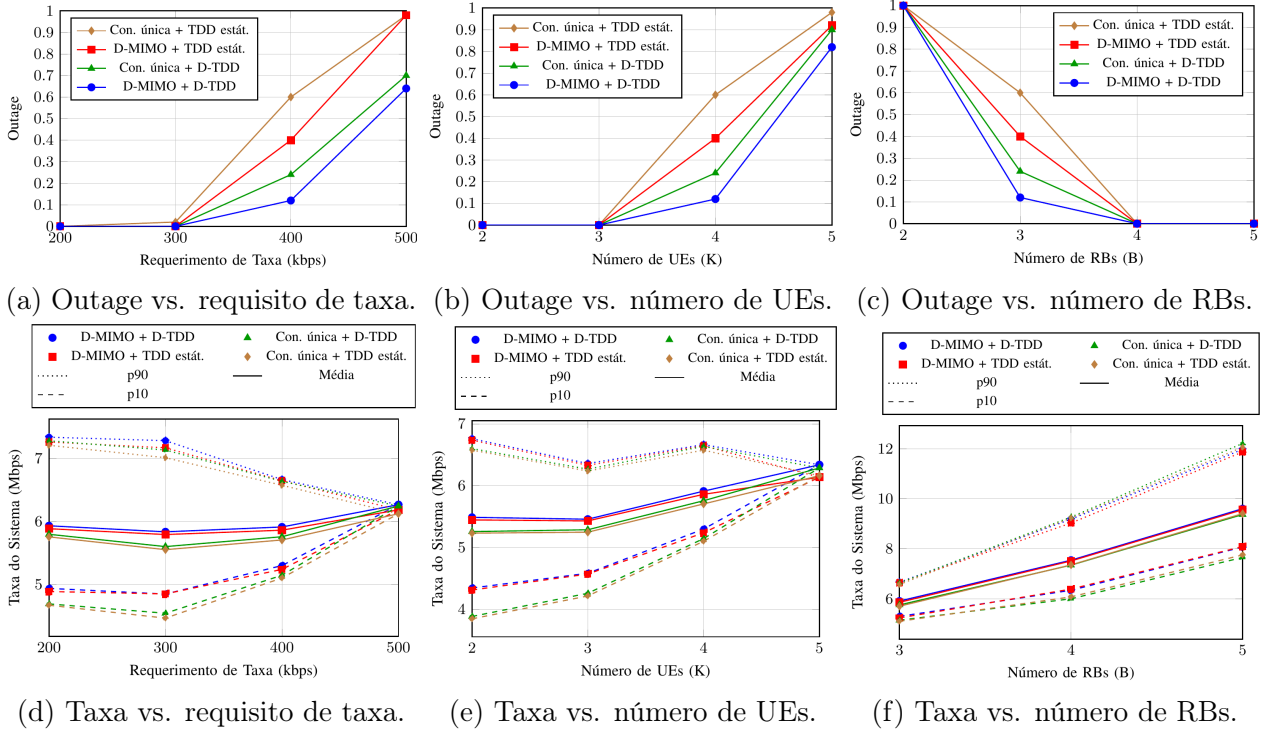


Figura 36: Resultados de desempenho para as arquiteturas avaliadas.

com D-MIMO e TDD estático. Embora a cooperação multi-AP aumente o ganho efetivo de canal por meio da combinação coerente em (110), o D-TDD amplia a flexibilidade de escalonamento ao permitir adaptação dinâmica das direções de transmissão em cada *slot*. Mesmo com requisitos simétricos em DL e UL, surgem naturalmente assimetrias temporais de tráfego devido às diferentes capacidades de transmissão e recepção dos nós da rede. Nesse contexto, a flexibilidade direcional proporcionada pelo D-TDD torna-se particularmente relevante em cenários com recursos limitados.

Com relação à taxa do sistema, Fig. 36d, observa-se redução da taxa agregada à medida que o requisito de QoS aumenta. Quando os requisitos são baixos, os recursos excedentes podem ser alocados a UEs com melhores condições de canal, aumentando a taxa total e ampliando a diferença entre os percentis de 10% e 90%. À medida que os requisitos mínimos crescem, mais recursos passam a ser utilizados para satisfazer as restrições de QoS, reduzindo a flexibilidade de alocação e aproximando as curvas correspondentes aos diferentes percentis.

B. Impacto do Número de UEs Figs. 36b e 36e apresentam a probabilidade de *outage* e a taxa do sistema em função do número de UEs, considerando $B = 3$ RBs e requisito mínimo de QoS de 400 kbps. À medida que K aumenta, a probabilidade de *outage* cresce para todas as arquiteturas. Com recursos totais fixos, o aumento do número de UEs reduz a quantidade média de recursos disponíveis por usuário, contraindo a região viável de (111).

Novamente, a operação centrada no usuário com D-MIMO e D-TDD apresenta o menor *outage*, enquanto conectividade única com D-TDD supera a operação centrada no usuário com D-MIMO e TDD estático. Esse resultado evidencia a importância da flexibilidade direcional. Sob TDD estático, toda a rede deve operar na mesma direção em cada *slot*, mesmo que apenas parte dos UEs necessite recursos adicionais em DL ou UL. Com o aumento de K , crescem também as assimetrias instantâneas de tráfego entre usuários, tornando o D-TDD mais eficiente

ao permitir decisões independentes de direção de transmissão para cada nó.

Com relação à taxa do sistema, comportamento semelhante ao observado anteriormente é verificado. Para pequenos valores de K , há recursos excedentes, favorecendo alocação oportunista e ampliando a diferença entre percentis. À medida que K aumenta, o sistema opera próximo ao limite da região viável, aproximando os diferentes percentis.

C. Impacto do Número de RBs Figs. 36c e 36f apresentam a probabilidade de *outage* e a taxa do sistema em função do número de RBs, considerando $K = 4$ UEs e requisito mínimo de QoS de 400 kbps. Observa-se redução do *outage* com o aumento de B , uma vez que mais RBs ampliam a região viável de (111). Para valores suficientemente elevados de B , o *outage* tende a zero em todas as arquiteturas, indicando que a abundância de recursos reduz o impacto das limitações estruturais das diferentes soluções.

Para valores moderados de B , mantém-se o mesmo ordenamento observado anteriormente: a operação centrada no usuário com D-MIMO e D-TDD apresenta o menor *outage*, enquanto conectividade única com D-TDD supera a operação centrada no usuário com D-MIMO e TDD estático. Isso confirma que, sob ortogonalização em frequência, a principal limitação de confiabilidade está associada à rigidez direcional do TDD estático, e não apenas aos ganhos de eficiência espectral proporcionados pela cooperação multi-AP.

Com relação à taxa do sistema, observa-se crescimento da taxa agregada com o aumento de B , consequência da disponibilidade de mais recursos para transmissão. Para pequenos valores de B , os percentis permanecem próximos, pois o sistema opera próximo ao limite mínimo viável. À medida que B aumenta, surgem mais recursos excedentes, permitindo maior alocação oportunista para UEs com melhores condições instantâneas de canal e ampliando a diferença entre os percentis de desempenho.

De forma geral, os resultados demonstram que a combinação entre operação centrada no usuário com D-MIMO e D-TDD proporciona a arquitetura mais robusta entre as soluções avaliadas. Enquanto a cooperação espacial aumenta os ganhos efetivos de canal, o D-TDD amplia a flexibilidade de escalonamento ao reduzir o acoplamento direcional imposto pelo TDD estático. A atuação conjunta desses mecanismos resulta em menor *outage* e maior taxa.

12.7 Conclusões e Trabalhos Futuros

Este trabalho investigou o gerenciamento de recursos de rádio em uma rede centrada no usuário com D-MIMO e operando com D-TDD. A formulação considerada realizou conjuntamente associação entre UEs e APs, alocação de potência e escalonamento de recursos tempo-frequência, resultando em um MINLP não convexo.

Para obter uma formulação tratável, o problema foi reformulado por meio da linearização das restrições de consistência TDD, ortogonalização em frequência entre UEs e parametrização convexa de potência baseada em *clusters*, permitindo sua reformulação como um MICP.

A avaliação de desempenho comparou arquiteturas com TDD estático e dinâmico combinadas com conectividade única e operação centrada no usuário com D-MIMO. Os resultados mostraram que a combinação entre operação centrada no usuário com D-MIMO e D-TDD proporciona menor *outage* e maior taxa do sistema. Além disso, observou-se que a flexibilidade direcional proporcionada pelo D-TDD pode ter impacto mais significativo na viabilidade do sistema do que a cooperação espacial isoladamente em cenários com recursos limitados.

Como perspectivas futuras, destacam-se estratégias distribuídas de otimização, técnicas avançadas de *beamforming* e abordagens não ortogonais de alocação de recursos para sistemas centrados no usuário com D-MIMO e D-TDD.

13 Agrupamento Centrado no Usuário e Controle de Potência em Redes Cell-Free com TDD Dinâmico

Matheus D. Carneiro, Samuel S. Silva,
Igor M. Guerreiro, Yuri C. B. Silva, Victor F. Monteiro
UFC

13.1 Introdução

As redes de comunicação sem fio têm evoluído para suportar um número crescente de usuários ativos, buscando atender a requisitos de maiores taxas de transmissão, menor latência e melhor cobertura de sinal. Nesse cenário, a arquitetura *Multiple-Input Multiple-Output* (MIMO) distribuída, também conhecida como *cell-free*, destaca-se como uma solução eficiente para as próximas gerações [116]. Caracterizada por um grande número de *Access Points* (APs) espalhados pela área de cobertura e processamento centralizado, essa estrutura permite a decodificação e o pré-processamento conjuntos, melhorando a qualidade do serviço.

Tradicionalmente, essas redes operam no modo *Time Division Duplex* (TDD) [116], onde a transmissão ocorre em uma única direção por vez, com uma divisão fixa entre *Uplink* (UL) e *Downlink* (DL). Entretanto, o crescimento heterogêneo das demandas de tráfego torna essa configuração estática ineficiente. Quando as necessidades dos usuários divergem, o TDD fixo degrada a qualidade do serviço e limita a capacidade do sistema de atender múltiplos dispositivos simultaneamente.

Para enfrentar esses desafios, a implementação do *Dynamic TDD* (D-TDD) tem sido estudada como uma possível solução para lidar com os problemas decorrentes da assimetria nas demandas de tráfego, permitindo um uso mais eficiente dos recursos temporais [125]. A implementação do D-TDD depende da arquitetura do sistema. Por exemplo, em arquiteturas celulares, as configurações de quadro podem ser ajustadas por célula para satisfazer às demandas dos usuários associados. Por outro lado, em sistemas *cell-free*, caracterizados por sua natureza centrada no usuário, os quadros de transmissão podem ser adaptados para cada usuário individualmente, sujeitos às restrições específicas da implantação da rede. Em redes com APs *half-duplex*, que só podem operar em um modo de transmissão por vez, uma implementação centrada no usuário é viável por meio do agrupamento centrado no usuário [125]. Essa abordagem garante que cada usuário seja atendido por um subconjunto disjunto ou exclusivo de APs, mitigando assim potenciais conflitos operacionais inerentes ao D-TDD.

Apesar das vantagens, o D-TDD introduz *Cross Link Interference* (CLI), que ocorre entre APs ou entre usuários que transmitem em sentidos opostos. Para lidar com esse desafio e aumentar a *Energy Efficiency* (EE), o controle de potência torna-se essencial. O trabalho em [133] propõe uma otimização conjunta do controle de potência e do ponto de chaveamento da direção de transmissão para cada *User Equipment* (UE) via aproximação convexa sucessiva para melhorar a EE. Abordagens clássicas de controle de potência, como o *Fractional Power Control* (FPC), também podem ser utilizadas. Por exemplo, [134] utiliza um algoritmo FPC em um cenário celular com D-TDD. O FPC é amplamente explorado na literatura com o objetivo de compensar parcialmente o efeito de *Large Scale Fading* (LSF).

Trabalhos anteriores analisaram o agrupamento centrado no usuário [135] e a alocação de potência [136] em um contexto *cell-free* D-TDD, mas como abordagens separadas e assumindo APs de antena única. Neste trabalho, estendemos esses estudos para analisar o impacto conjunto de ambas as técnicas, além de generalizar o cenário para suportar APs multiantena.

Partindo inicialmente de um agrupamento disjunto, considera-se o uso de *Dynamic Cooperation Clustering* (DCC) para aumentar o tamanho dos grupos, ao mesmo tempo em que são satisfeitas as restrições de direção dos enlaces. Em seguida considera-se a implementação de um esquema FPC em adição ao DCC, sendo avaliado o desempenho do sistema em termos de métricas como a eficiência espectral, eficiência energética e atraso de pacotes.

13.2 Modelo de Sistema

Neste trabalho, consideramos uma rede MIMO *cell-free* operando sob TDD dinâmico, a qual é composta por L APs e K UEs. Os UEs estão distribuídos aleatoriamente pela área de cobertura, enquanto os APs estão dispostos em uma malha regular. Os UEs são equipados com uma única antena omnidirecional, ao passo que os APs possuem um arranjo de N antenas. Os APs estão conectados a uma única CPU por meio de enlaces de *fronthaul* e estão sujeitos à restrição *half-duplex*, na qual a transmissão e a recepção não podem ocorrer simultaneamente.

13.2.1 Modelo de Propagação

Os canais de propagação — incluindo os enlaces UE-para-AP, UE-para-UE e AP-para-AP — seguem uma distribuição de Rice, consistindo em componentes *Line of Sight* (LOS) e componentes difusas. O canal de propagação $\mathbf{H}_{k,l} \in \mathbb{C}^{N_t \times N_r}$ entre o k -ésimo transmissor e o l -ésimo receptor pode ser expresso como:

$$\mathbf{H}_{k,l} = \left[\sqrt{\frac{\bar{K}_{k,l}}{\bar{K}_{k,l} + 1}} \mathbf{A}_{k,l} + \sqrt{\frac{1}{\bar{K}_{k,l} + 1}} \mathbf{H}^{\omega}_{k,l} \right] \beta_{k,l}, \quad (114)$$

em que $\bar{K}_{k,l}$ denota o fator K de Rice; $\mathbf{H}^{\omega}_{k,l} \sim (\mathbf{0}, \mathbf{R}_{k,l})$ é a componente *Non-Line of Sight* (NLOS), modelada seguindo uma distribuição Rayleigh correlacionada; $\mathbf{A}_{k,l} \in \mathbb{C}^{N_t \times N_r}$ é a matriz de direcionamento; $\beta_{k,l}$ é o coeficiente LSF, que modela os efeitos de perda de percurso e sombreamento.

A probabilidade de LOS é inerente a todos os enlaces do sistema e depende da posição dos elementos no ambiente e de suas alturas. A probabilidade de LOS pode ser definida pela seguinte expressão:

$$\text{Prob}_{\text{LOS}}(d) = 0.5 - \min(0.5, 5e^{-0.156/d}) + \min(0.5, e^{-d/0.03}), \quad (115)$$

onde d é a distância entre o transmissor e o receptor em km. A determinação do efeito de perda de percurso, a geração do efeito de sombreamento e o fator K de Rice dependem diretamente se o canal é LOS ou não. Se o canal for LOS, $\bar{K} \sim \mathcal{N}(9,5)$ dB; caso contrário, $\bar{K}$ é igual a zero. Consideramos que o enlace entre dois UEs não possui componente de percurso direto, portanto, $\bar{K}$ é sempre igual a zero.

O coeficiente de LSF $\beta_{k,l}$ é definido como $\beta_{k,l} = 10^{\frac{PL_{k,l} + SH_{k,l}}{10}}$, onde $SH_{k,l} \sim \mathcal{N}(0, \sigma_{sh}^2)$ e $PL_{k,l}$ representam os ganhos de sombreamento espacialmente correlacionado e de perda de percurso em dB. O cálculo da perda de percurso depende do tipo de enlace sendo modelado e das condições de LOS de acordo com [132]; o valor de σ_{sh}^2 depende unicamente das condições de LOS. O cálculo da perda de percurso para todos os enlaces e condições de LOS é implementado como em [132] e a geração do sombreamento é implementada como em [137].

13.2.2 Combinadores

Após o cálculo das estimativas de canal entre os UEs e os APs, os coeficientes serão utilizados para configurar os combinadores de recepção e os precodificadores de transmissão. Para mitigar a interferência multiusuário e a interferência cruzada, o conjunto de APs $\mathcal{A}_k$ que atende ao UE k no UL aplica um combinador *Minimum Mean Square Error* (MMSE). Por outro lado, para direcionar melhor os feixes em direção ao UE i , cada AP no conjunto $\mathcal{A}_i$ aplicará um precodificador *Maximum Ratio* (MR). Além disso, denotamos o conjunto de UEs no UL como $\mathcal{U}$ e o conjunto de UEs no DL como $\mathcal{D}$.

O combinador MMSE aplicado pelo conjunto de APs $\mathcal{A}_k$ é dado por:

$$\mathbf{v}_k = p_k \left(\sum_{i \in \mathcal{U}} p_i \mathbf{\Psi}_{i, \mathcal{A}_k} + \sum_{j \in \mathcal{D}} \sum_{l \in \mathcal{A}_j} \mathbf{Q}_{l, \mathcal{A}_k} \right)^{-1} \hat{\mathbf{h}}_{k, \mathcal{A}_k}, \quad (116)$$

em que $\mathbf{\Psi}_{i, \mathcal{A}_k} = (\mathbf{h}_{i, \mathcal{A}_k} \mathbf{h}_{i, \mathcal{A}_k}^H + \mathbf{C}_i)$ é a componente de interferência multiusuário, sendo $\mathbf{h}_{i, \mathcal{A}_k}$ o canal coletivo entre o UE i e todos os APs em $\mathcal{A}_k$, e $\mathbf{C}_{i, \mathcal{A}_k} = \text{diag}(\mathbf{C}_{i,1}, \dots, \mathbf{C}_{i,|\mathcal{A}_k|})$ é a matriz diagonal da correlação do erro de estimação de canal, conforme [116]; $\mathbf{Q}_{l, \mathcal{A}_k}$ representa a componente de interferência entre APs, que pode ser expressa como:

$$\mathbf{Q}_{l, \mathcal{A}_k} = \mathbf{G}_{l, \mathcal{A}_k} \mathbf{p}_{l,j} \mathbf{p}_{l,j}^H \mathbf{G}_{l, \mathcal{A}_k}^H, \quad (117)$$

onde $\mathbf{G}_{l, \mathcal{A}_k} \in \mathbb{C}^{N|\mathcal{A}_k| \times N}$ denota a matriz de canal coletiva do AP interferente l para o agrupamento de APs $\mathcal{A}_k$, e $\mathbf{p}_{l,j} \in \mathbb{C}^N$ é o precodificador no DL que o AP interferente l utiliza para transmitir para o UE j .

13.2.3 Modelo de Sinal

Esta seção apresenta as expressões de sinal e de *Spectral Efficiency* (SE) para uma rede *cell-free* que utiliza D-TDD. Como resultado da comutação dinâmica entre as direções de transmissão, os níveis de interferência variam e as SEs de todos os UEs precisam ser recalculadas constantemente. Denotamos o conjunto de UEs transmitindo ativamente em UL e DL durante um determinado *slot* de tempo n como $\mathcal{U}$ e $\mathcal{D}$.

O sinal recebido pelo conjunto de APs que atende ao UE k no UL pode ser expresso como:

$$\mathbf{y}_{\mathcal{A}_k} = \sqrt{p_k} \mathbf{h}_{k, \mathcal{A}_k} x_k + \gamma_{\mathcal{A}_k}, \quad (118)$$

em que $\mathbf{h}_{k, \mathcal{A}_k} \in \mathbb{C}^{N|\mathcal{A}_k|}$ é o vetor de canal coletivo para o UE k de todos os APs em $\mathcal{A}_k$; $p_k \in \mathbb{R}$ e x_k são a potência de transmissão e o símbolo do k -ésimo UE, respectivamente. O primeiro termo da expressão é o sinal desejado e o segundo, γ , é a soma da interferência e do ruído, sendo dado por:

$$\gamma_{\mathcal{A}_k} = \sum_{\substack{i \in \mathcal{U} \\ i \neq k}} \sqrt{p_i} \mathbf{h}_{i, \mathcal{A}_k} x_i + \sum_{\substack{j \in \mathcal{D} \\ j \neq k}} \sum_{l \in \mathcal{A}_j} \mathbf{G}_{l, \mathcal{A}_k} \mathbf{p}_{l,j} + \mathbf{n}_{\mathcal{A}_k},$$

onde $\mathbf{h}_{i, \mathcal{A}_k} \in \mathbb{C}^{N|\mathcal{A}_k|}$ é o vetor de canal coletivo entre o i -ésimo UE interferente e todos os APs em $\mathcal{A}_k$; $p_i \in \mathbb{R}$ e $x_i \in \mathbb{C}$ são a potência de transmissão e o símbolo do i -ésimo UE; $\mathbf{G}_{l, \mathcal{A}_k} \in \mathbb{C}^{N|\mathcal{A}_k| \times N}$ é o canal coletivo entre o AP de DL l e o agrupamento de APs $\mathcal{A}_k$, $\mathbf{p}_{l,j} \in \mathbb{C}^N$ é o precodificador de DL que o AP l usa para transmitir para o j -ésimo UE e $\mathbf{n}_{\mathcal{A}_k} \in \mathbb{C}^N$ é o ruído de recepção. A SE do UE k em UL é dada por:

$$R_k^{\text{ul}} = \log_2 (1 + \text{SINR}_k^{\text{ul}}), \quad (119)$$

onde a SINR_k pode ser calculada como:

$$\text{SINR}_k^{\text{ul}} = \frac{p_k |\mathbf{v}_k^H \mathbf{h}_{k, \mathcal{A}_k} x_k|^2}{|\mathbf{v}_k^H \boldsymbol{\gamma}_{\mathcal{A}_k}|^2}. \quad (120)$$

O sinal recebido pelo k -ésimo UE em DL é dado por:

$$y_k = \sum_{a \in \mathcal{A}_k} \mathbf{h}_{a,k}^H \mathbf{p}_{a,k} s_k + \psi_k, \quad (121)$$

onde o primeiro termo é a superposição dos sinais desejados, com $\mathbf{h}_{a,k} \in \mathbb{C}^N$ e $\mathbf{p}_{a,k} \in \mathbb{C}^N$ sendo o vetor de canal e o precodificador entre o k -ésimo UE e o AP a , que pertence a $\mathcal{A}_k$, e ψ_k é a soma das componentes de interferência e ruído:

$$\psi_k = \sum_{\substack{i \in \mathcal{D} \\ i \neq k}} \sum_{l \in \mathcal{A}_i} \mathbf{h}_{l,k}^H \mathbf{p}_{l,i} s_i + \sum_{j \in \mathcal{U}} \sqrt{p_j} h_{j,k} s_j + n_k, \quad (122)$$

na qual o primeiro termo é a interferência multiusuário de DL e o segundo é a componente de interferência entre UEs, sendo $\mathbf{h}_{l,k} \in \mathbb{C}^N$ o canal entre o AP de DL l e o UE vítima k , $h_{j,k} \in \mathbb{C}$ o canal para o enlace entre o j -ésimo e o k -ésimo UEs e $n_k \in \mathbb{C}$ o ruído de recepção em DL. A SE do UE k em DL é dada por:

$$R_k^{\text{dl}} = \log_2 (1 + \text{SINR}_k^{\text{dl}}), \quad (123)$$

onde a $\text{SINR}_k^{\text{dl}}$ pode ser calculada como:

$$\text{SINR}_k^{\text{dl}} = \frac{|\mathbf{h}_{k, \mathcal{A}_k}^H \mathbf{p}_{\mathcal{A}_k} s_k|^2}{|\psi_k|^2}. \quad (124)$$

Para avaliar o desempenho do sistema, a EE por usuário é definida como a razão entre a eficiência espectral alcançada e a potência de transmissão consumida. Para o UL, a EE do k -ésimo UE é dada por:

$$\text{EE}_k^{\text{ul}} = \frac{R_k^{\text{ul}}}{p_k}. \quad (125)$$

Para o DL, a EE contabiliza a potência total alocada ao k -ésimo UE pelo seu agrupamento de APs, $\mathcal{A}_k$:

$$\text{EE}_k^{\text{dl}} = \frac{R_k^{\text{dl}}}{\sum_{l \in \mathcal{A}_k} p_{l,k}} \quad (126)$$

onde $p_{l,k}$ é a potência de transmissão em DL alocada pelo l -ésimo AP ao k -ésimo UE.

13.3 Controle de Potência

Para mitigar a CLI e melhorar a EE, adotamos uma abordagem amplamente explorada para definir os níveis de potência dos UEs no UL, o FPC. Para fins de comparação, adotamos a *Max Power* (MP) como estratégia básica, na qual cada UE transmite com a potência máxima $p_k = P_{\max} \forall k \in \{1, \dots, K\}$ no modo UL.

Considerando o esquema FPC, aplicado anteriormente a sistemas MIMO massivo [138] e *cell-free* [139, 140], o seu objetivo é compensar parcialmente o efeito de perda de percurso de acordo com um fator de compensação α e uma constante multiplicativa P_0 . Assim, a potência

de transmissão de um UE k é dada por $p_k = \min(P_{\max}, P_0 \zeta_k^{-\alpha})$, no qual ζ_k considera os coeficientes de LSF e o vetor de direcionamento dos APs associados a cada UE k conforme:

$$\zeta_k = \sqrt{\sum_{a \in \mathcal{A}_k} \text{tr} \left(\frac{\beta_{k,a}}{\bar{K}_{k,a} + 1} [\bar{K}_{k,a} \mathbf{A}_{k,a} \mathbf{A}_{k,a}^H + \mathbf{I}_{N_L}] \right)}. \quad (127)$$

Note que a regra de alocação de potência baseada em (127) visa primordialmente manter um nível de desempenho alvo para cada UE (escalonado por P_0) enquanto reduz o gasto energético total. Consequentemente, ela não contabiliza explicitamente os termos de interferência, além de limitar implicitamente a potência de transmissão individual em $P_{\max}$.

13.4 Estratégia de Agrupamento em D-TDD

Neste trabalho, propomos uma estratégia de agrupamento primária projetada para acomodar as configurações de tráfego específicas de cada UE individualmente. Para garantir que o serviço seja adaptado às suas necessidades, impomos uma restrição em que cada AP se conecta apenas a UEs operando no mesmo modo de transmissão. Inicialmente, alocamos um número fixo e igual de APs para cada UE com base nos coeficientes de LSF para formar agrupamentos disjuntos. Sob essa alocação inicial, os UEs com os maiores coeficientes de LSF são atribuídos aos melhores APs disponíveis. Consequentemente, alguns UEs podem ser relegados a APs com ganhos de canal comparativamente menores.

Para lidar com essa limitação e aprimorar o desempenho geral do sistema, propomos um algoritmo de realocação que permite que UEs com baixos ganhos de canal sejam atendidos por APs superiores, desde que esses APs já estejam atendendo a outros UEs na mesma direção de transmissão. Esse processo, denominado "re-agrupamento", é inerentemente dinâmico devido à variabilidade do tráfego e das direções de transmissão. No entanto, ele permanece computacionalmente viável na prática porque as configurações de quadro são definidas antecipadamente, permitindo que a *Central Processing Unit* (CPU) calcule a alocação dinâmica de APs antes da transmissão.

Implementamos esse re-agrupamento adaptando o DCC, um algoritmo amplamente estudado na arquitetura D-MIMO. Embora o DCC em [116] tenha sido originalmente projetado para cenários de TDD estático, que não exigem um agrupamento disjunto inicial, nós o adaptamos para funcionar como uma abordagem de re-agrupamento.

Este algoritmo realoca os APs para cada UE com base na SNR. O primeiro passo desta abordagem consiste em identificar, para cada UE k , o AP no agrupamento inicial com a melhor SNR, definindo-o como o AP mestre com $\text{SNR}_{\text{master}}$. Em seguida, avaliamos todos os APs que não estão atualmente no agrupamento e alocamos aqueles que satisfazem a condição $\text{SNR}_a - \text{SNR}_{\text{master}} \geq \delta, \forall a \in \{1, \dots, L\} \setminus \mathcal{A}_k$ no modo UL, com ambos os valores de SNR em dB. O limiar δ determina o quão restritivo o algoritmo pode ser ao adicionar um novo AP.

13.5 Resultados de Simulação

A rede *cell-free* simulada com D-TDD consiste em $K = 25$ UEs e $L = 100$ APs distribuídos em uma área de cobertura de 1.2 km^2 . A combinação MMSE é considerada para na recepção no UL e o processamento MR é considerado na transmissão no DL. A potência total do AP é dividida igualmente entre os fluxos dos UEs que ele atende. Parâmetros de simulação adicionais podem ser encontrados na Tabela 11.

Tabela 11: Parâmetros de simulação da rede *cell-free* com D-TDD

Parâmetro	Valor
Frequência da portadora	$f_c = 2$ GHz
Altura dos equipamentos (AP / UE)	11.5 / 1.5 m
Número de pilotos	$\tau_p = 10$
Comprimento do bloco de coerência	$\tau_c = 200$
Potência máx. de transmissão (UL / DL)	0.1 W / 0.2 W
FPC P_0 e α	-10 dBm e 0.5
Densidade de potência de ruído	-174 dBm/Hz
Variância do sombreamento (LOS, NLOS)	4, 8.2
Duração do slot	1 ms
Tamanho dos pacotes UL e DL	50 kbytes
Intervalo de chegada de pacotes	(UL, DL)
UL-heavy (predomínio de UL)	(8,16) ms
DL-heavy (predomínio de DL)	(16,8) ms
Simétrico	(12,12) ms

O sistema de tráfego de pacotes baseia-se em um modelo de filas em tempo discreto, com a chegada de pacotes modelada por um processo de Poisson, e com o enfileiramento e desenfileiramento ocorrendo exclusivamente no início e no final de cada slot de tempo. Além disso, o intervalo médio entre a chegada de pacotes varia de acordo com as demandas individuais dos usuários, sendo menor para direções de tráfego com maior demanda. O atraso do pacote é definido como o número de slots necessários para que um determinado pacote seja totalmente transmitido, abrangendo o intervalo desde o momento em que entra na fila até a sua saída.

13.5.1 Análise do Limiar DCC

Primeiramente, analisamos o efeito do algoritmo de reagrupamento baseado em DCC no desempenho do sistema. Para esta análise, apresentamos curvas de atraso e SE tanto em UL quanto em DL, de acordo com o limiar adotado $\delta \in \{-15, -10, 0, 10, 20, 30, 40\}$ dB, tendo como algoritmo de referência o agrupamento disjunto. Aqui, consideramos o MP para a realização do controle de potência.

Nesta configuração, a Figura 37 apresenta o número médio de APs conectados a cada UE tanto em UL quanto em DL. As primeiras barras mostram o caso em que nenhuma reclusterização é realizada e os APs são alocados igualmente entre os UEs; as demais mostram os casos com DCC em ordem decrescente de δ . Pelos resultados, observa-se que, conforme δ aumenta, o algoritmo realoca menos APs para cada UE. O caso com $\delta = 40$ dB apresenta o mesmo valor médio de APs realocados que o caso disjunto. Já para $\delta = -15$ dB, a reclusterização é a menos restritiva, adicionando, em média, mais de 5 APs para o modo UL e mais de 8 APs para o modo DL.

A Figura 38a ilustra os valores de atraso nos percentis 10, 50 e 90 variando com diferentes limiares, bem como os casos disjuntos de referência. Aqui, podemos observar claramente que, a partir de $\delta = -15$ dB, o atraso diminui continuamente até atingir um ponto mínimo em $\delta = 10$ dB e, em seguida, começa a aumentar. Isso é mais perceptível no 90º percentil, indicando que, no UL, o DCC é mais sensível ao limiar para o UE com o pior desempenho de atraso.

A Figura 38b ilustra os percentis 10, 50 e 90 da SE no UL como uma função de δ , juntamente

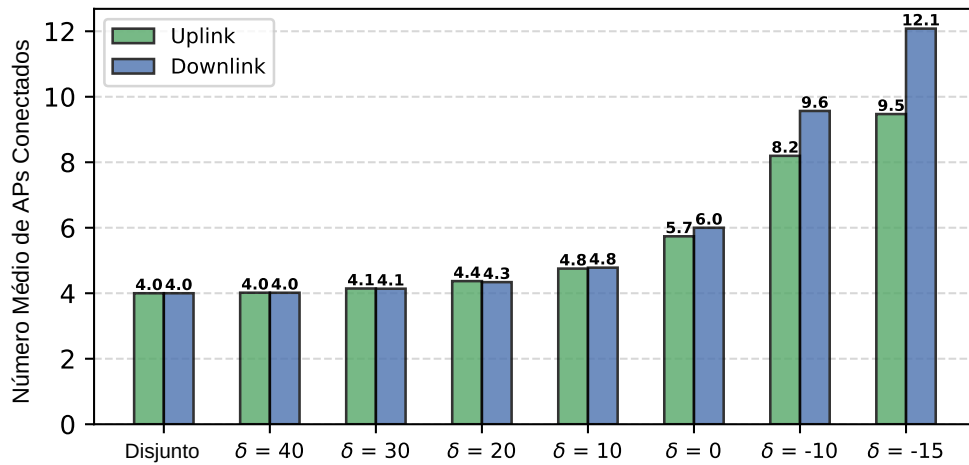
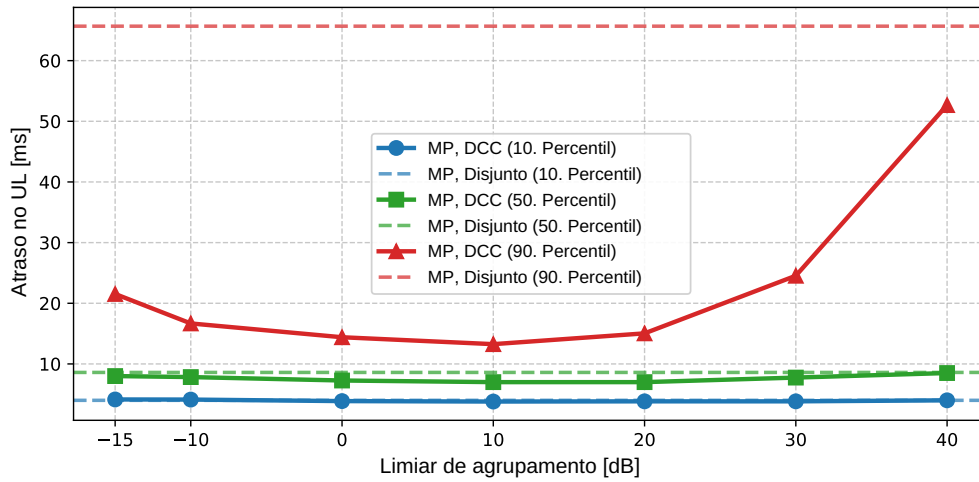
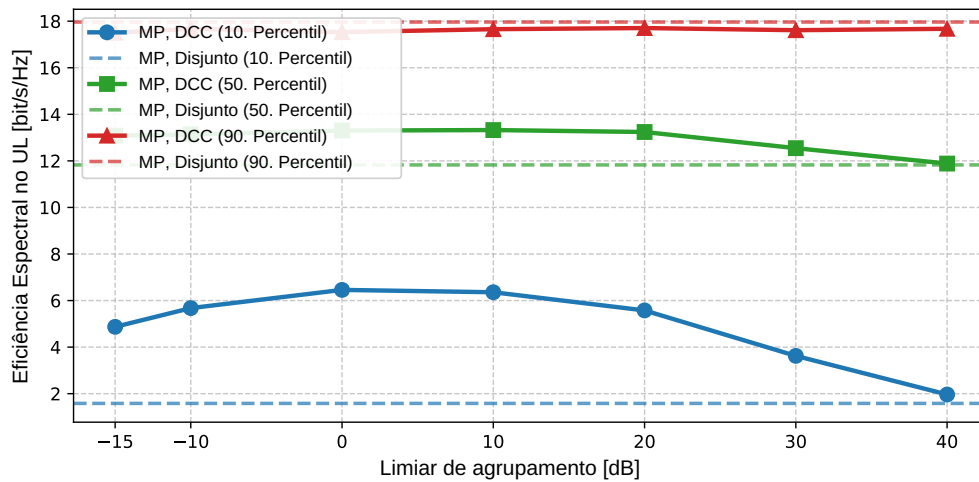


Figura 37: Número médio de APs conectados a cada UE para valores de δ .

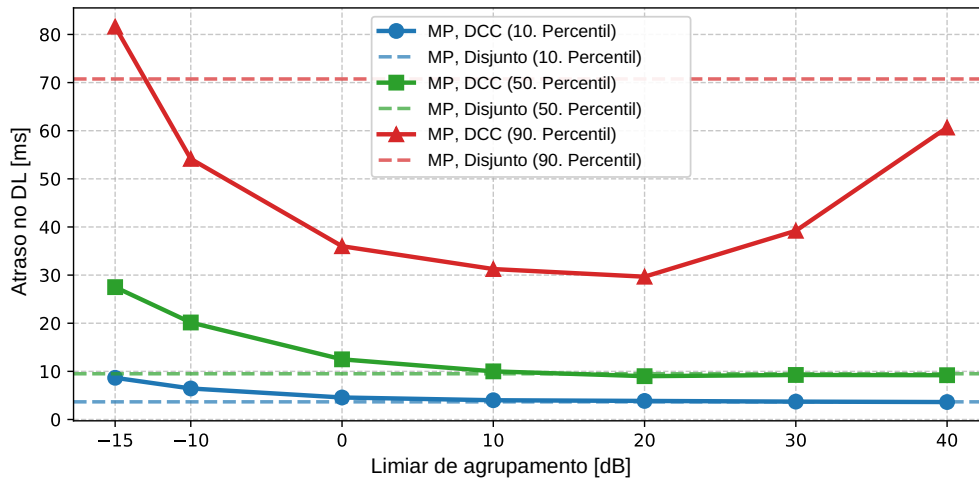


(a) Atrasos em UL nos percentis 10, 50 e 90.

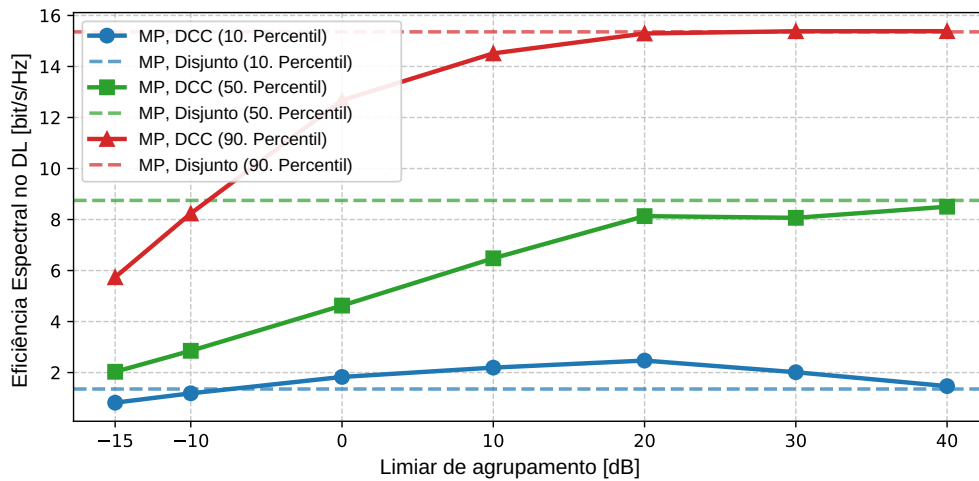


(b) SE em UL nos percentis 10, 50 e 90.

Figura 38: Impacto do limiar do DCC (δ) nas métricas de desempenho em UL.



(a) Atrasos em DL nos percentis 10, 50 e 90.



(b) SE em DL nos percentis 10, 50 e 90.

Figura 39: Impacto do limiar do DCC (δ) nas métricas de desempenho em DL.

com os casos disjuntos. Observamos que, para os percentis 90 e 50, o impacto do δ é marginal, uma vez que esses usuários já desfrutavam de condições de canal robustas. Em contraste, o 10º percentil mostra uma sensibilidade significativa ao limiar; a SE aumenta até atingir um máximo entre $\delta = 0$ dB e $\delta = 10$ dB. Além deste ponto, o desempenho declina em direção ao caso disjunto à medida que δ se torna mais restritivo. É evidente que os ganhos proporcionados pelo DCC nos percentis mais baixos de SE traduzem-se diretamente em melhorias nos percentis mais altos de atraso. Isso ocorre porque o agrupamento disjunto inicial prioriza UEs com melhores condições de canal, enquanto UEs desfavorecidos são inicialmente atribuídos a APs com ganhos menores. Reduzir o valor de δ permite que esses UEs sejam realocados para APs de maior qualidade, promovendo maior uniformidade espacial em todo o sistema.

Passando agora para a análise de DL, a Figura 39a ilustra os percentis 10, 50 e 90 do atraso em DL como uma função de δ , juntamente com os casos disjuntos. Para os percentis 10 e 50, o atraso diminui conforme δ aumenta até eventualmente convergir para o caso disjunto a partir de $\delta = 10$ dB. Para o 90º percentil, o atraso diminui até um valor mínimo em $\delta = 20$ dB antes de aumentar novamente. Ao associar mais UEs a um AP específico em valores mais

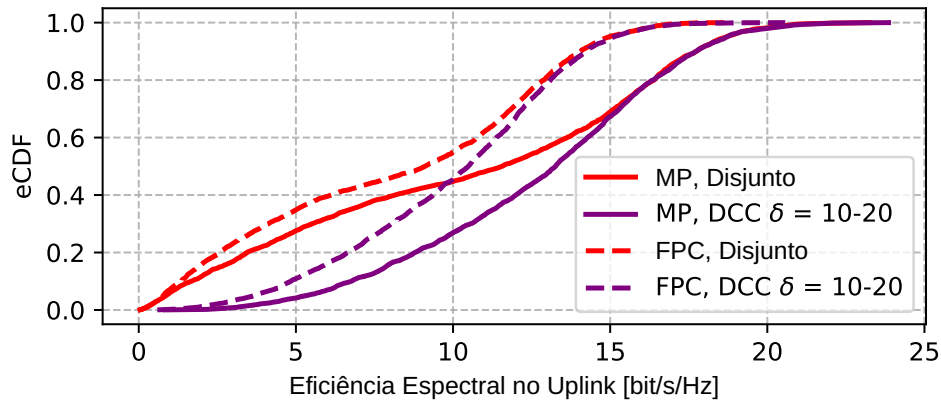
baixos de δ , a transmissão em DL sofre interferência amplificada e divisão de potência, o que consequentemente aumenta o atraso de transmissão. À medida que δ se torna mais restritivo, essa interferência excessiva é mitigada.

A Figura 39b ilustra os percentis 10, 50 e 90 da SE em DL como uma função de δ , juntamente com os casos disjuntos. Conforme o limiar aumenta, os valores de 50º e 90º aumentam até tenderem ao caso disjunto. Começando de $\delta = 40$ dB no 50º e $\delta = 20$ dB no 90º. Para o 10º percentil, o valor sobe até o ponto máximo em $\delta = 20$ dB e começa a diminuir novamente.

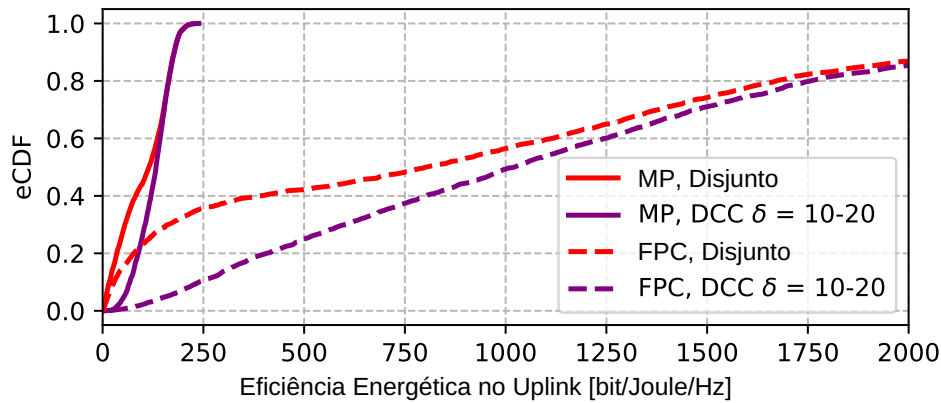
Esses resultados sugerem que os melhores limiares de DCC são: $\delta = 10$ dB para o UL e $\delta = 20$ dB para o DL, pois apresentam o melhor desempenho geral entre as métricas analisadas.

13.5.2 Análise do Controle de Potência e Agrupamento

Em seguida, analisamos o efeito do algoritmo de reagrupamento DCC e do algoritmo de controle de potência FPC no desempenho do sistema. Para esta análise, apresentamos as curvas de SE e EE na direção de UL (visto que o DL não é significativamente afetado pelo FPC) em quatro configurações: i) MP com agrupamento disjunto ii) MP com DCC iii) FPC com agrupamento disjunto iv) FPC com DCC. Conforme estabelecido anteriormente, para o DCC, definimos os limiares como $\delta = 10$ dB para UL e $\delta = 20$ dB para DL. Para o FPC, utilizamos a constante multiplicativa $P_0 = -10$ dBm e o fator de compensação $\alpha = 0.5$.



(a) eCDF da UL SE.



(b) eCDF da UL EE.

Figura 40: Comparação de esquemas de controle de potência e agrupamento nas métricas de desempenho de UL.

A Figura 40a ilustra a *Empirical Cumulative Distribution Function* (eCDF) da SE em UL para esta configuração. Comparando os esquemas de agrupamento, também podemos observar um comportamento consistente com nossas descobertas iniciais. Em relação às estratégias de controle de potência dentro de cada esquema de agrupamento, o MP supera o FPC em todos os percentis. Isso ocorre porque o FPC não incorpora explicitamente a mitigação de interferência e, portanto, reduz seu próprio nível de potência em relação ao MP para garantir um desempenho suficiente, o que diminui a taxa de transmissão de dados. No entanto, o caso com FPC e DCC é superior ao caso com MP e agrupamento disjunto nos percentis mais baixos, ocorrendo o oposto nos percentis mais altos.

A Figura 40b ilustra a eCDF da EE em UL para esta configuração. Para a EE em UL, a comparação entre os esquemas de agrupamento é a mesma da SE, na qual o DCC supera o agrupamento disjunto, especialmente para percentis mais baixos. Quanto aos esquemas de controle de potência, é evidente que o FPC proporciona melhor EE do que o MP. Isso acontece porque, no FPC, os níveis de potência não podem exceder a potência máxima e geralmente apresentam valores menores. Assim, a redução do denominador na Equação (125) é mais significativa do que a redução no numerador quando comparado ao MP, o que melhora consideravelmente a razão.

13.6 Conclusões

Neste trabalho, avaliamos os efeitos conjuntos dos algoritmos de reagrupamento dinâmico (DCC) e de controle de potência (FPC) no atraso de pacotes, SE e EE dentro de uma rede MIMO *cell-free* com D-TDD, comparando-os com os algoritmos de agrupamento disjunto e MP.

Nossas descobertas iniciais indicam que, sob valores específicos de limiar, o DCC melhora significativamente a SE e, conseqüentemente, reduz o atraso de pacotes tanto nas direções de UL quanto de DL. No UL, o DCC supera o algoritmo disjunto em todos os casos avaliados. Essa melhoria é mais pronunciada para UEs desfavorecidos, pois o algoritmo permite que eles sejam atendidos por APs de maior qualidade fora de seus agrupamentos iniciais. Por outro lado, no DL, o DCC apenas supera o caso disjunto dentro de uma faixa específica de valores de limiar. Configurações altamente permissivas introduzem interferência excessiva tanto no UL quanto no DL devido aos novos enlaces de transmissão criados pelas alocações adicionais de APs, mas esse efeito é mais forte no DL. Embora esse mecanismo auxilie UEs desfavorecidos, ele pode degradar severamente o desempenho de usuários que já possuem condições de canal robustas.

Em relação à análise conjunta de controle de potência e agrupamento, os resultados indicam que o MP supera o FPC em termos de UL SE na maioria dos percentis. Uma exceção notável é observada para UEs desfavorecidos, onde o FPC combinado com o DCC resulta em uma SE melhor do que o caso de MP com agrupamento disjunto. Por outro lado, o FPC supera vastamente o MP em termos de EE em todas as configurações. Isso torna o FPC uma estratégia altamente atrativa para a otimização de recursos energéticos. Finalmente, examinando os resultados de DL, o FPC supera marginalmente o MP em termos de SE e EE. Essa melhoria entre enlaces cruzados ocorre primordialmente porque a potência de transmissão reduzida no UL, inerente ao FPC, introduz menos *Inter-User Interference* (IUI) aos UEs que estão recebendo no DL.

14 Estratégias de Alocação de Potência para Mensagens Comuns e Privadas em Redes RSMA LEO-Terrestres

Michel Bernardo de Paiva, Francisco Rafael Marques Lima
CPQD

14.1 Introdução

A crescente demanda por comunicações sem fio com alta capacidade e baixa latência tem impulsionado o desenvolvimento de técnicas avançadas de acesso múltiplo [141]. Nesse contexto, o *Rate-Splitting Multiple Access* (RSMA) surge como uma solução promissora, devido à sua capacidade de gerenciar interferências de forma eficiente e aumentar a eficiência espectral [142]. Ao explorar o princípio de divisão de taxa, o RSMA possibilita estratégias de comunicação mais flexíveis e robustas quando comparadas aos esquemas tradicionais, como *Orthogonal Multiple Access* (OMA), *Non-Orthogonal Multiple Access* (NOMA) e *Space Division Multiple Access* (SDMA) [142].

Paralelamente, a implantação de redes de satélites em órbita baixa da Terra (*Low Earth Orbit* (LEO)) tem ganhado destaque como uma alternativa para fornecer cobertura global e conectividade de alta velocidade [143]. A integração do RSMA com sistemas de comunicação via satélite LEO representa uma oportunidade relevante para otimizar a alocação de recursos, aumentar a confiabilidade do sistema e maximizar a taxa de transmissão total (*throughput*) [144]. No entanto, alcançar um desempenho ótimo nesses sistemas integrados exige o uso de técnicas avançadas de otimização.

Em [145], os autores apresentam uma análise comparativa da taxa de dados obtida com RSMA, SDMA e OMA em um sistema *downlink* multi-feixe via satélite LEO. Considerando precodificadores convencionais para fluxos comum e privados, e um cenário simplificado com dois usuários, a taxa de dados do RSMA é otimizada por meio de uma busca exaustiva na divisão de potência entre as mensagens comum e privadas. Os resultados de simulação demonstram que o RSMA é capaz de alcançar desempenho superior em termos de taxa de transmissão quando comparado às abordagens tradicionais.

Os resultados de simulação apresentados demonstram que o RSMA é capaz de alcançar taxas de dados superiores ou equivalentes às obtidas por SDMA e OMA. Em [146], os autores implementam um modelo de sistema semelhante ao de [145], porém com o objetivo de avaliar uma solução de pré-codificação baseada em aprendizado de máquina (*Machine Learning* – ML) aplicada ao SDMA. Os resultados indicam que a solução proposta apresenta desempenho superior, especialmente em cenários com informação imperfeita do estado do canal no transmissor (*Channel State Information at the Transmitter* – CSIT).

Neste trabalho, considera-se a otimização da taxa total de dados em sistemas de comunicação baseados em satélites LEO utilizando RSMA. Para isso, são empregados esquemas de pré-codificação para as mensagens comum e privadas, bem como uma estratégia de alocação de potência entre essas mensagens. A abordagem de alocação de potência adotada apresenta menor complexidade computacional em comparação com [145], que se baseia em busca exaustiva.

14.2 Modelo do Sistema e formulação do problema

Nesta seção, é apresentado o modelo do sistema considerado, bem como a formulação do problema de maximização da taxa total de dados em um sistema baseado em RSMA aplicado

a comunicações via satélite LEO.

14.2.1 Modelo do Sistema

Considera-se um sistema de comunicação *downlink* via satélite LEO utilizando RSMA em um canal com visada direta (*Line-of-Sight* – LOS). O satélite LEO é equipado com um arranjo linear uniforme de antenas (*Uniform Linear Array* – ULA), composto por N elementos, cada um com ganho G_{sat} . O sistema atende K usuários, cada um equipado com uma única antena e ganho G_{user} .

O ângulo de saída do sinal transmitido pelo arranjo em direção ao usuário k é representado por β_k . A menor distância entre o satélite e a superfície terrestre é denotada por d_{sat} . Os usuários são distribuídos uniformemente ao redor do centro da célula, localizado a uma distância D do nadir do satélite. O canal LOS entre o satélite e o usuário k é representado pelo vetor $\mathbf{h}_k \in \mathbb{C}^{N \times 1}$.

O sinal recebido pelo usuário k pode ser expresso por:

$$y_k = \mathbf{h}_k^H \mathbf{x} + n_k = \sum_{n=1}^N h_{k,n} x_n + n_k, \quad (128)$$

em que $\mathbf{x} \in \mathbb{C}^{N \times 1}$ representa o sinal pré-codificado transmitido pelo satélite, x_n é o sinal transmitido pela n -ésima antena, $h_{k,n}$ é a resposta do canal entre a antena n e o usuário k , e $n_k \sim \mathcal{CN}(0, \sigma_n^2)$ representa o *Additive White Gaussian Noise* (AWGN) complexo.

A resposta do canal entre a n -ésima antena do satélite e o usuário k é dada por:

$$h_{k,n}(\cos(\beta_k)) = \sqrt{G_{user} G_{sat}} \frac{\lambda_{carr}}{4\pi d_k} e^{-j\omega_k} a_{k,n}(\cos(\beta_k)), \quad (129)$$

onde d_k é a distância entre o satélite e o usuário k , λ_{carr} é o comprimento de onda da portadora, $\omega_k \in [0, 2\pi]$ representa o deslocamento total de fase, e $a_{k,n}(\cos(\beta_k))$ corresponde ao deslocamento de fase relativo associado ao arranjo de antenas.

O termo de fase relativo é definido por:

$$a_{k,n}(\cos(\beta_k)) = \exp \left(-j\pi \frac{d_a}{\lambda_{carr}} (N + 1 - 2n) \cos(\beta_k) \right), \quad (130)$$

em que d_a representa o espaçamento entre os elementos do arranjo ULA.

Definindo o vetor de direcionamento do arranjo como:

$$\mathbf{a}_k(\cos(\beta_k)) = [a_{k,1}(\cos(\beta_k)), \dots, a_{k,N}(\cos(\beta_k))]^T, \quad (131)$$

o vetor de canal pode ser escrito como:

$$\mathbf{h}_k(\cos(\beta_k)) = \sqrt{G_{user} G_{sat}} \frac{\lambda_{carr}}{4\pi d_k} e^{-j\omega_k} \mathbf{a}_k(\cos(\beta_k)). \quad (132)$$

Devido aos altos desvios Doppler presentes em sistemas LEO, considera-se informação imperfeita do estado do canal no transmissor (*Channel State Information at the Transmitter* – CSIT). Essa imperfeição é modelada por uma perturbação no cosseno do ângulo de saída, dada por $\tau_k \sim \mathcal{U}(-\tau, +\tau)$, resultando em:

$$\mathbf{h}'_k(\cos(\beta_k)) = \mathbf{h}_k(\cos(\beta_k)) \circ \mathbf{a}_k(\tau_k), \quad (133)$$

onde $\circ$ representa o produto de Hadamard. Essa modelagem indica que erros na estimativa do ângulo de saída afetam diretamente o vetor de canal utilizado no projeto dos pré-codificadores e, conseqüentemente, o desempenho do sistema RSMA.

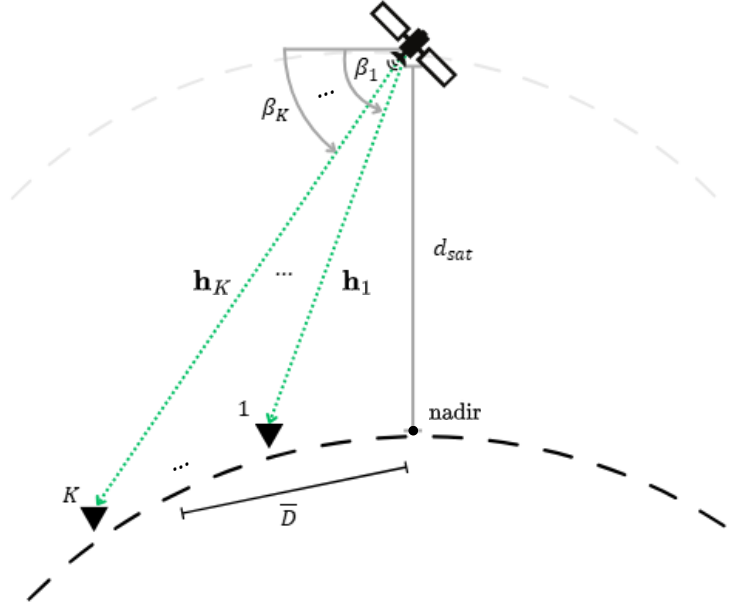


Figura 41: Modelo do sistema de satélite LEO.

14.2.2 Modelo de Canal e Informação de Estado

O modelo de canal adotado considera propagação LOS, onde a resposta entre a n -ésima antena e o usuário k depende da distância, do ganho das antenas e do deslocamento de fase. Devido às altas velocidades envolvidas em sistemas LEO, há a presença de efeitos Doppler significativos.

Assim, assume-se que a informação de estado do canal no transmissor (*Channel State Information at the Transmitter* (CSIT)) é imperfeita, sendo modelada por uma perturbação no cosseno do ângulo de saída. Essa imprecisão impacta diretamente o projeto dos pré-codificadores e, conseqüentemente, o desempenho do sistema [145].

14.3 Modelo de Transmissão com RSMA

Para a aplicação do RSMA, a mensagem de cada usuário W_k é dividida em duas partes: uma parte comum ($W_{c,k}$) e uma parte privada ($W_{p,k}$). As partes comuns de todos os usuários são combinadas em uma única mensagem comum global, que é codificada em um fluxo comum s_c e transmitida a todos os usuários.

As partes privadas são codificadas individualmente em fluxos distintos $s_1, \dots, s_K$. O sinal transmitido pode ser expresso como:

$$\mathbf{x} = \mathbf{p}_c s_c + \sum_{k=1}^K \mathbf{p}_k s_k \quad (134)$$

onde $\mathbf{p}_c$ e $\mathbf{p}_k$ são os vetores de pré-codificação para o fluxo comum e os fluxos privados, respectivamente.

No receptor, cada usuário realiza inicialmente a decodificação do fluxo comum, tratando os sinais privados como interferência. Em seguida, aplica-se cancelamento sucessivo de interferência (*Successive Interference Cancellation* – SIC), permitindo a remoção do sinal comum e a posterior decodificação do fluxo privado correspondente.

A taxa associada ao fluxo comum para o usuário k é dada por:

$$R_{c,k} = \log_2 \left(1 + \frac{|\mathbf{h}_k^H \mathbf{p}_c|^2}{\sigma^2 + \sum_{i=1}^K |\mathbf{h}_k^H \mathbf{p}_i|^2} \right) \quad (135)$$

Para garantir que todos os usuários consigam decodificar a mensagem comum, a taxa comum do sistema é definida como:

$$R_c = \min_{k \in \mathcal{K}} R_{c,k} \quad (136)$$

A taxa do fluxo privado do usuário k é expressa por:

$$R_{p,k} = \log_2 \left(1 + \frac{|\mathbf{h}_k^H \mathbf{p}_k|^2}{\sigma^2 + \sum_{i \neq k} |\mathbf{h}_k^H \mathbf{p}_i|^2} \right) \quad (137)$$

A taxa total alcançável para o usuário k é dada pela soma das taxas comum e privada:

$$R_k = R_c + R_{p,k} \quad (138)$$

14.3.1 Formulação do Problema de Otimização

O objetivo do sistema é maximizar a taxa total de dados, definida como a soma das taxas individuais dos usuários. Para isso, são consideradas como variáveis de otimização os vetores de pré-codificação e a alocação de potência entre os fluxos comum e privados.

O problema pode ser formulado como:

$$\max_{\mathbf{p}, s} \sum_{k=1}^K R_k \quad (139)$$

sujeito à restrição de potência total:

$$\sum_{k \in \mathcal{K}} \|\mathbf{p}_k\|^2 + \|\mathbf{p}_c\|^2 \leq P_{max} \quad (140)$$

onde P_{max} representa a potência total disponível para transmissão.

Esse problema é, em geral, não convexo, o que torna necessária a adoção de estratégias de solução com menor complexidade computacional, conforme discutido nas seções seguintes.

14.4 Solução Proposta

O problema formulado apresenta elevada complexidade, uma vez que tanto a função objetivo quanto as restrições envolvidas são não convexas em relação às variáveis de otimização. Neste trabalho, utiliza-se uma solução de baixa complexidade baseada nas abordagens propostas em [147] e [148], nas quais o método *Zero Forcing* (ZF) é empregado para as mensagens privadas, enquanto um precodificador do tipo *max-min data rate* é utilizado para a mensagem comum. Além disso, a análise é limitada a um cenário simplificado com dois usuários.

14.4.1 Análise dos Precodificadores

A aplicação do ZF consiste em fixar a direção dos precodificadores e ajustar a potência das transmissões, de modo que cada usuário receba seu sinal desejado com máxima intensidade, enquanto os sinais destinados aos demais usuários apresentem interferência desprezível.

O vetor de precodificação ZF do usuário k é definido como

$$\mathbf{p}_k = \sqrt{P_k} \frac{\mathbf{c}_k}{\|\mathbf{c}_k\|}, \quad (141)$$

em que $\mathbf{c}_k$ é obtido a partir da pseudo-inversa da matriz de canal $\mathbf{H}$:

$$\mathbf{P}_{ZF} = \mathbf{H}(\mathbf{H}^H \mathbf{H})^{-1}, \quad (142)$$

sendo $\mathbf{H} = [\mathbf{h}_1 \dots \mathbf{h}_K]$, $\mathbf{c}_k$ a k -ésima coluna de $\mathbf{P}_{ZF}$ e P_k a potência alocada ao fluxo privado do usuário k .

Dessa forma, obtém-se $|\mathbf{h}_k^H \mathbf{p}_i| = 0$, para $i \neq k$, e

$$|\mathbf{h}_k^H \mathbf{p}_k|^2 = \|\mathbf{h}_k\|^2 \rho P_k, \quad (143)$$

em que $\rho \in [0,1]$ representa a correlação entre os canais, sendo $\rho = 0$ para canais completamente alinhados e $\rho = 1$ para canais ortogonais.

O fator de correlação é definido a partir do determinante da matriz de Gram, composta pelos produtos internos entre os vetores de canal normalizados. Para $K = 2$, tem-se

$$\rho = 1 - |\hat{\mathbf{h}}_1^H \hat{\mathbf{h}}_2|^2, \quad (144)$$

onde $\hat{\mathbf{h}}_k = \frac{\mathbf{h}_k}{\|\mathbf{h}_k\|}$.

Por outro lado, o precodificador da mensagem comum é projetado como o vetor de *beamforming* capaz de maximizar a taxa comum. Assim, define-se o seguinte problema de otimização:

$$\max_{\mathbf{p}_c} \min \left(\frac{|\mathbf{h}_1^H \mathbf{p}_c|^2}{\sigma^2 + |\mathbf{h}_1^H \mathbf{p}_1|^2}, \frac{|\mathbf{h}_2^H \mathbf{p}_c|^2}{\sigma^2 + |\mathbf{h}_2^H \mathbf{p}_2|^2} \right) \quad (145)$$

sujeito a

$$\|\mathbf{p}_c\|^2 = P_c. \quad (146)$$

Definindo $\psi_k^2 = \sigma^2 + |\mathbf{h}_k^H \mathbf{p}_k|^2 = \sigma^2 + \|\mathbf{h}_k\|^2 \rho P_k$ e $\tilde{\mathbf{h}}_k = \frac{\mathbf{h}_k}{\psi_k}$, o problema pode ser reescrito como

$$\max_{\mathbf{p}_c} \min \left(|\tilde{\mathbf{h}}_1^H \mathbf{p}_c|^2, |\tilde{\mathbf{h}}_2^H \mathbf{p}_c|^2 \right) \quad (147)$$

sujeito a

$$\|\mathbf{p}_c\|^2 = P_c. \quad (148)$$

A solução desse problema é dada por

$$\mathbf{p}_c = \sqrt{P_c} \mathbf{f}_c, \quad (149)$$

onde $\mathbf{f}_c$ representa a direção do precodificador, obtida por

$$\mathbf{f}_c = \frac{1}{\sqrt{\lambda}} \left(\mu_1 \tilde{\mathbf{h}}_1 + \mu_2 \tilde{\mathbf{h}}_2 e^{-j\angle\alpha_{12}} \right), \quad (150)$$

com

$$\lambda = \frac{\alpha_{11}\alpha_{22} - |\alpha_{12}|^2}{\alpha_{11} + \alpha_{22} - 2|\alpha_{12}|}, \quad (151)$$

$$\begin{bmatrix} \mu_1 \\ \mu_2 \end{bmatrix} = \frac{1}{\alpha_{11} + \alpha_{22} - 2|\alpha_{12}|} \begin{bmatrix} \alpha_{22} - |\alpha_{12}| \\ \alpha_{11} - |\alpha_{12}| \end{bmatrix}, \quad (152)$$

sendo

$$\begin{bmatrix} \alpha_{11} & \alpha_{12} \\ \alpha_{12}^* & \alpha_{22} \end{bmatrix} = \begin{bmatrix} \tilde{\mathbf{h}}_1^H \\ \tilde{\mathbf{h}}_2^H \end{bmatrix} \begin{bmatrix} \tilde{\mathbf{h}}_1 & \tilde{\mathbf{h}}_2 \end{bmatrix}. \quad (153)$$

Além disso, define-se

$$r_k^2 = 1 + \frac{|\mathbf{h}_k^H \mathbf{p}_k|^2}{\sigma^2} = 1 + \frac{\|\mathbf{h}_k\|^2 \rho P_k}{\sigma^2}. \quad (154)$$

Assim, a taxa total obtida pelos usuários utilizando RSMA é dada por

$$R_{sum} = R_c + \log_2(r_1^2) + \log_2(r_2^2), \quad (155)$$

em que R_c satisfaz o problema anteriormente formulado.

Como $|\tilde{\mathbf{h}}_1^H \mathbf{p}_c| = |\tilde{\mathbf{h}}_2^H \mathbf{p}_c|$, a taxa comum pode ser escrita como

$$R_c = \log_2 \left(1 + |\tilde{\mathbf{h}}_2^H \mathbf{p}_c|^2 \right). \quad (156)$$

Além disso, considerando $\psi_k^2 = \sigma^2 r_k^2$, a taxa total torna-se

$$R_{sum} = \log_2(r_1^2) + \log_2 \left(r_2^2 + \frac{|\mathbf{h}_2^H \mathbf{p}_c|^2}{\sigma^2} \right). \quad (157)$$

14.4.2 Divisão de Potência

Sejam P_{max} a potência total disponível para transmissão, P_p a potência destinada às mensagens privadas e P_c a potência destinada à mensagem comum, tais que $P_p = sP_{max}$ e $P_c = (1 - s)P_{max}$.

A potência entre os fluxos privados é otimizada utilizando a estratégia de *water-filling* (*Water-Filling* (WF)). Dessa forma, a potência destinada ao fluxo privado do usuário k é dada por

$$P_k = \frac{1}{K} \left(\sum_{\forall i} \frac{\sigma^2}{\rho \|\mathbf{h}_i\|^2} + P_p \right) - \frac{\sigma^2}{\rho \|\mathbf{h}_k\|^2}. \quad (158)$$

Seguindo [147], define-se

$$\Gamma = \frac{1}{\rho} \left(\frac{1}{\|\mathbf{h}_2\|^2} - \frac{1}{\|\mathbf{h}_1\|^2} \right), \quad (159)$$

que representa simultaneamente a correlação entre os canais e a diferença de intensidade entre eles.

Assim, são considerados dois regimes principais de operação do RSMA.

1) OMA/NOMA/Multicast: Quando $sP_{max} \leq \Gamma$, define-se $P_2 = 0$ e $P_1 = sP_{max}$. Nesse regime, o RSMA comporta-se como:

- Multicast, para $s = 0$;
- NOMA, para $0 < s < 1$;
- OMA, para $s = 1$.

Assim, o parâmetro s é ajustado de modo que o RSMA opere no esquema de acesso múltiplo mais eficiente para esse regime.

2) SDMA/RSMA: Por outro lado, quando $sP_{max} > \Gamma$, os valores de P_1 e P_2 são determinados pela solução de WF, sendo ambos maiores que zero.

Nesse caso, o RSMA comporta-se como SDMA quando $s = 1$, enquanto para $0 < s < 1$ o sistema opera em seu modo híbrido característico.

Substituindo r_k^2 , ψ_k^2 e P_k em (157), obtém-se

$$R_{sum} = \log_2 \left[\frac{ac + (ad + bc)s + bds^2}{(\sigma^2)^2} \right], \quad (160)$$

onde

$$b = \frac{\|\mathbf{h}_1\|^2 \rho P_{max}}{2}, \quad a = \sigma^2 + \frac{\Gamma}{P_{max}} b,$$

$$d = \frac{\|\mathbf{h}_2\|^2 \rho P_{max}}{2} - |\mathbf{h}_2^H \mathbf{f}_c|^2 P_{max}, \quad (161)$$

e

$$c = \sigma^2 - \frac{\Gamma}{P_{max}} d + |\mathbf{h}_2^H \mathbf{f}_c|^2 (P_{max} - \Gamma). \quad (162)$$

Além disso, quando $P_1 > 0$ e $P_2 > 0$, a expressão da taxa total deixa de depender diretamente de s e da diferença entre as intensidades dos canais, passando a depender apenas das direções dos canais [147].

Assim, a solução ótima é obtida impondo

$$\frac{\partial R_{sum}}{\partial s} = 0, \quad (163)$$

o que resulta em $s = -\frac{a}{2b} - \frac{c}{2d}$.

Portanto, o valor ótimo de s para esse regime é dado por

$$s = \min \left(-\frac{a}{2b} - \frac{c}{2d}, 1 \right). \quad (164)$$

14.5 Avaliação de Desempenho

Nesta seção, são apresentados os detalhes de implementação, as soluções de referência e os resultados obtidos, com o objetivo de analisar as vantagens e limitações da abordagem proposta. A Tabela 12 resume os principais parâmetros utilizados nas simulações, adotados com base em [145].

Tabela 12: Parâmetros de simulação da rede LEO-terrestre.

Parâmetro	Variável	Valor
Altitude do satélite	d_{sat}	600 km
Frequência da portadora	f	2 GHz
Tamanho do arranjo do satélite	N	6
Espaçamento entre antenas	d_a	7.5 cm
Ganho da antena do satélite	G_{sat}	16 dBi
Ganho da antena do usuário	G_{user}	0 dBi
Potência de ruído	σ_n^2	-122 dBW
Potência máxima de transmissão do satélite	P_{max}	25 dBW

14.5.1 Detalhes de Implementação

Foi considerado um cenário simplificado composto por dois usuários, distribuídos uniformemente ao redor de um ponto central da célula localizado a uma distância D do nadir do satélite. Dessa forma, a distância do usuário k ao nadir do satélite é modelada como uma variável aleatória $D_k \sim \mathcal{U}(D - \Delta D, D + \Delta D)$. Além disso, é garantida uma separação mínima entre os usuários, também igual a ΔD .

Para avaliar o desempenho das técnicas de acesso múltiplo, dois cenários principais são analisados:

- **Cenário (i):** Variação da distância D , considerando $\Delta D = 10$ km, com dois casos de informação de estado do canal no transmissor (CSIT):
 - $\tau = 0$, representando CSIT perfeito;
 - $\tau = 0,05$, representando CSIT imperfeito.
- **Cenário (ii):** Distância fixa $D = 100$ km e $\Delta D = 10$ km, com variação do nível de imperfeição do CSIT, representado pelo parâmetro τ .

Em ambos os cenários, são realizadas 1000 repetições de Monte Carlo, com o objetivo de garantir a confiabilidade estatística dos resultados obtidos. Os demais parâmetros de simulação seguem as configurações adotadas em [145] e trabalhos correlatos da literatura.

14.5.2 Soluções de Referência

Para fins de comparação, foram consideradas soluções de referência baseadas em RSMA, OMA e SDMA, conforme apresentado em [145]. No trabalho de referência, a alocação de potência em RSMA é definida a partir do parâmetro α , otimizado por busca exaustiva. Assim, a potência destinada ao fluxo comum é dada por $P_c = P_{max} - P_{max}^\alpha$, enquanto a potência

destinada aos fluxos privados é dada por $P_p = P_{\max}^\alpha$, em que $P_{\max}$ representa a potência total disponível no sistema.

Para o fluxo comum, é utilizado um precodificador fixo, definido por

$$\mathbf{p}_c = \sqrt{\frac{P_c}{N}} \mathbf{1}.$$

Já para os fluxos privados, é adotado o precodificador baseado no critério de mínimo erro quadrático médio, *Minimum Mean Square Error* (MMSE), cuja matriz de pré-codificação é expressa por

$$\mathbf{P}_p = \sqrt{\frac{P_p}{\text{tr}(\mathbf{W}^H \mathbf{W})}} \mathbf{W},$$

em que

$$\mathbf{W} = \left(\mathbf{H}^H \mathbf{H} + \frac{\sigma_n^2 K}{P_{\max}} \mathbf{I}_N \right)^{-1} \mathbf{H}^H.$$

A técnica MMSE também é empregada na implementação de SDMA; entretanto, nesse caso, a matriz de pré-codificação é escalonada pela potência total disponível, uma vez que não há divisão das mensagens em partes comum e privada. Por fim, para a implementação de OMA, utiliza-se o precodificador de máxima razão de transmissão, *Maximal Ratio Transmission* (MRT), definido para o k -ésimo usuário como

$$\mathbf{p}_k^{\text{MRT}} = \sqrt{\frac{P_{\max}}{K}} \frac{\mathbf{h}_k^H}{\|\mathbf{h}_k\|}.$$

Essas soluções são utilizadas como referência para avaliar o desempenho da abordagem proposta, permitindo comparar a taxa total alcançável e a robustez frente às imperfeições de informação de canal no transmissor.

14.6 Resultados

Nesta seção, apresentamos uma avaliação detalhada do esquema proposto de RSMA, comparando seu desempenho com técnicas convencionais de acesso múltiplo, a saber, OMA, SDMA e uma implementação de referência de RSMA.

1) Impacto da Distância dos Usuários (D)

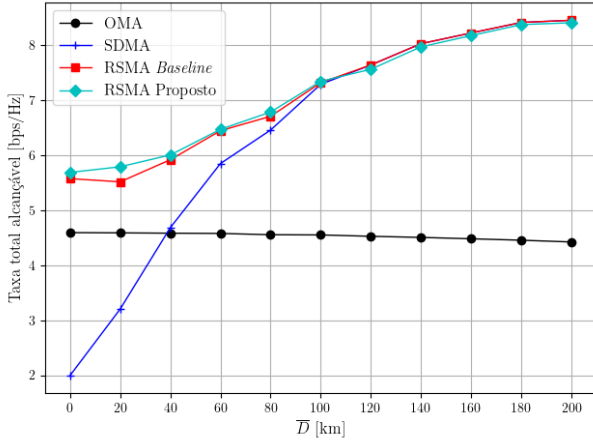
Inicialmente, analisamos o impacto da distância dos usuários em relação ao satélite. A Fig. 42b apresenta os resultados sob condições de CSIT perfeito.

Observa-se que o OMA apresenta comportamento relativamente estável à medida que D aumenta, uma vez que elimina a interferência entre usuários por meio da alocação ortogonal de recursos, ainda que com perda em eficiência espectral.

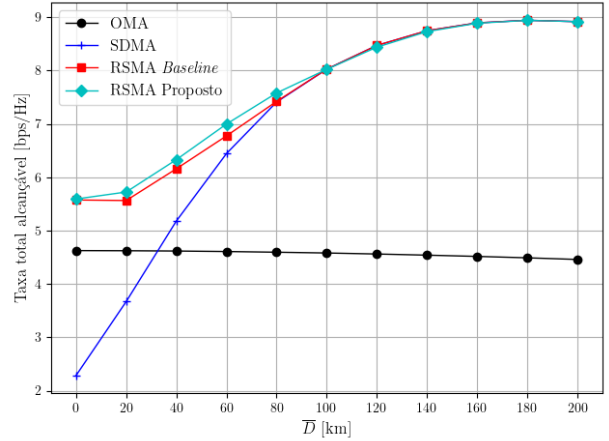
Por outro lado, SDMA e RSMA apresentam ganhos significativos com o aumento de D , devido à melhor separação espacial entre os usuários, o que favorece as estratégias de pré-codificação. Em particular, o SDMA explora multiplexação espacial, enquanto o RSMA combina decodificação de interferência com mitigação parcial da mesma.

Nota-se ainda que o RSMA proposto alcança taxas iguais ou superiores às do RSMA de referência em toda a faixa de distâncias analisadas, com ganhos mais expressivos em regiões próximas ao satélite.

A Fig. 42a apresenta os resultados considerando CSIT imperfeito. Nesse cenário, todas as técnicas sofrem degradação de desempenho devido a erros na estimação do canal.



(a) CSIT Imperfeito



(b) CSIT Perfeito

Figura 42: Taxa total de dados alcançada pelos usuários em função de D , para avaliar o desempenho de *Multiple Access* (MA). A simulação foi conduzida com $\Delta D = 10$ km, considerando $\tau = 0$ para CSIT perfeito e $\tau = 0.05$ para CSIT imperfeito.

O SDMA é particularmente afetado, evidenciando sua alta dependência de informações precisas de canal. O RSMA de referência mantém maior robustez, enquanto o método proposto supera o benchmark para distâncias de até aproximadamente 120 km. Após esse ponto, observa-se uma leve degradação no desempenho, embora não significativa, indicando que a solução proposta ainda é eficaz em uma ampla faixa de operação.

2) Impacto da Imperfeição do CSIT (τ)

A Fig. 43 apresenta o comportamento da taxa total alcançável em função do nível de imperfeição do CSIT, representado por τ .

À medida que τ aumenta, a precisão da informação de estado de canal diminui, impactando diretamente o desempenho das técnicas de pré-codificação. O SDMA apresenta degradação acentuada, reforçando sua sensibilidade a erros de estimação.

O OMA foi omitido nesta análise devido ao seu desempenho inferior. Em contrapartida, o RSMA apresenta uma degradação mais gradual, resultado de sua capacidade de dividir as mensagens em partes comum e privada, permitindo tratar parte da interferência via decodificação.

O método proposto de RSMA supera consistentemente o RSMA de referência para todos os valores de τ , conforme ilustrado na Fig. 43, demonstrando maior robustez frente a imperfeições no CSIT.

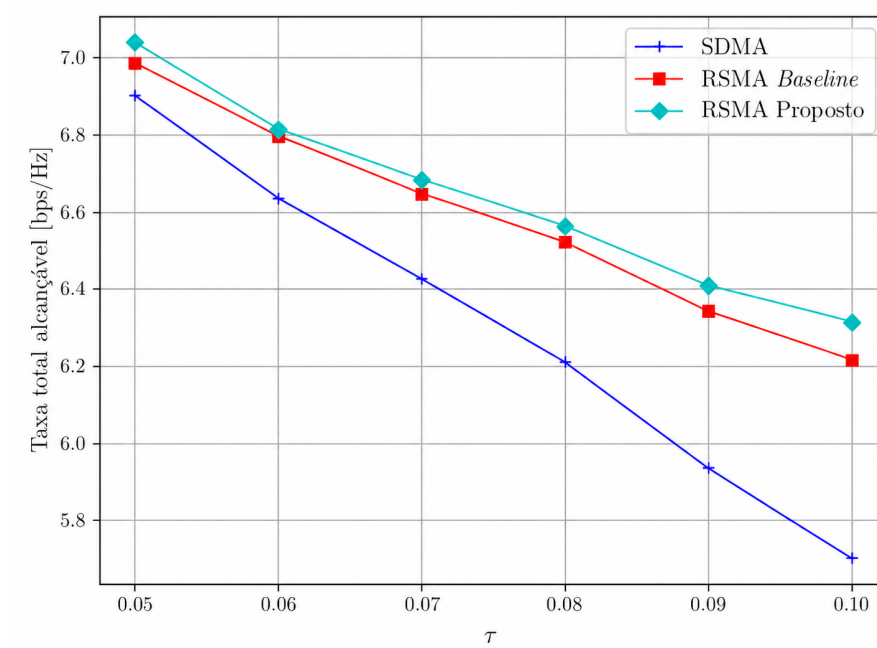


Figura 43: Taxa total de dados alcançada pelos usuários em função de τ , para avaliar o desempenho de MA. A simulação foi conduzida com $\Delta D = 10$ km e $D = 100$ km.

14.7 Conclusão

Este estudo investigou como a distribuição de potência e o desenho dos pré-codificadores influenciam o desempenho do RSMA em cenários de comunicação via satélite. A partir da formulação de um problema de maximização de taxa, foi desenvolvida uma estratégia de baixa complexidade que combina um pré-codificador comum orientado ao critério max-min com a técnica ZF para as mensagens privadas. Essa escolha permitiu não apenas simplificar a análise teórica, mas também viabilizar uma otimização eficiente da potência, resultando em um sistema mais adaptável e resiliente, capaz de superar abordagens convencionais mesmo sob condições adversas de canal, conforme evidenciado pelos resultados de simulação.

Como continuidade, há espaço para ampliar o modelo considerando um número maior de usuários e os efeitos de imperfeições no processo de *Successive Interference Cancellation* (SIC). Também se mostra promissora a incorporação de requisitos de qualidade de serviço e o uso de técnicas de otimização mais avançadas, incluindo abordagens clássicas e baseadas em aprendizado de máquina, com potencial para tornar a solução ainda mais aplicável em cenários reais.

15 Avaliação de RSMA Hierárquico em Sistemas Downlink com Múltiplas Antenas: Outage e Alocação de Potência

Raphael Parreira de Souza e Edgar Eduardo Benitez Olivo
CPqD e UNICAMP

15.1 Introdução

As redes sem fio de próxima geração devem suportar conectividade massiva sob restrições rigorosas de eficiência espectral e confiabilidade. Nesse contexto, o *Rate-Splitting Multiple Access* (RSMA) surge como uma solução flexível para o problema central de gerenciamento da interferência multiusuário. Ao dividir a mensagem de cada usuário em componente comum e componente privada, o RSMA unifica e generaliza dois paradigmas tradicionais: o tratamento da interferência como ruído, característico do *Space Division Multiple Access* (SDMA), e sua decodificação completa, como no *Non-Orthogonal Multiple Access* (NOMA) [149].

Com base no esquema convencional de camada única, o *Hierarchical Rate-Splitting Multiple Access* (HRSMA) amplia as capacidades de eficiência espectral e mitigação de interferência em cenários onde os usuários podem ser agrupados [149, 150]. Nessa abordagem, cada mensagem é decomposta em três partes: comum ao sistema, comum ao grupo e privada. Isso permite gerenciar simultaneamente a interferência entre grupos, por meio do fluxo comum ao sistema, e a interferência intra-grupo, por meio dos fluxos comuns de grupo. O modelo mantém elevada flexibilidade, uma vez que a anulação das taxas comuns de grupo reduz o HRSMA ao RSMA convencional. Entretanto, essa estrutura em múltiplas camadas aumenta a complexidade do *Successive Interference Cancellation* (SIC), exigindo a decodificação sequencial das camadas comuns ao sistema e ao grupo.

Embora a maior parte dos trabalhos recentes se concentre na robustez do RSMA convencional [151, 152, 153, 154], há um crescente interesse na investigação do potencial do HRSMA [150, 155, 156]. Um estudo de referência apresenta uma análise abrangente do desempenho do RSMA em enlace descendente multiusuário sob desvanecimento Nakagami- m [151]. Nesse trabalho, são derivadas expressões exatas e assintóticas para a probabilidade de indisponibilidade (*Outage Probability* (OP)), taxa de erro de símbolo e capacidade ergódica, além da otimização da alocação de potência para minimizar a OP e maximizar a taxa ergódica soma. No entanto, a análise limita-se ao esquema de camada única.

A arquitetura hierárquica foi posteriormente introduzida em [157], com o uso de agrupamento de usuários via k -means, baseado na correlação de canal. Esse modelo também avalia a taxa soma assintótica sob desvanecimento Rayleigh espacialmente correlacionado e com *Channel State Information at the Transmitter* (CSIT) imperfeito. Trabalhos mais recentes investigam o HRSMA em conjunto com tecnologias assistivas. Em [155], os usuários são agrupados considerando similaridade e intensidade de canal, permitindo transmissão assistida por *Reconfigurable Intelligent Surfaces* (RIS) e avaliação de desempenho via taxa soma sob desvanecimento Rician. Em [156], uma metasuperfície inteligente empilhada é otimizada por meio do ajuste de fase, alocação de potência e agrupamento dinâmico de usuários, visando maximizar a equidade do sistema sem o uso de pré-codificação digital.

Apesar desses avanços, ainda existem lacunas importantes quanto aos limites teóricos e ao desempenho máximo da arquitetura hierárquica. Essas questões tornam-se especialmente relevantes em redes subcarregadas, onde o número de antenas transmissoras é maior ou igual ao número de usuários. Nesses cenários, os graus de liberdade espaciais podem ser explorados para

simplificar a detecção dos fluxos privados. Técnicas de pré-codificação linear, como *Zero Forcing* (ZF) combinada com *Maximal Ratio Transmission* (MRT), permitem eliminar completamente a interferência multiusuário nos fluxos privados [151].

A análise da OP nessas condições fornece uma base fundamental para compreender o desempenho do HRSMA e sustentar sua aplicação em sistemas sem fio altamente confiáveis. Além disso, o impacto do número de grupos e das estratégias de agrupamento de usuários, considerando diferentes condições de canal, ainda não é completamente compreendido.

Neste trabalho, é desenvolvido um arcabouço teórico abrangente para análise e otimização do HRSMA em enlace descendente sob desvanecimento Nakagami- m , estabelecendo um benchmark de desempenho para sistemas com pré-codificação digital. As principais contribuições são estas:

- Obtenção de expressões fechadas, exatas e assintóticas, para a OP no HRSMA, incorporando a interferência intra-grupo inerente à transmissão hierárquica;
- Formulação de um problema de otimização do tipo min-max para minimizar a pior OP, com proposta de um algoritmo eficiente baseado em Progressão Geométrica (PG), considerando pré-codificação ZF-MRT em cenários subcarregados;
- Validação dos resultados analíticos por meio de simulações de Monte Carlo, além da análise do impacto de parâmetros críticos do sistema, como configuração de antenas, heterogeneidade de canal e estratégias de agrupamento de usuários.

15.2 Modelo do Sistema

Considera-se um sistema de comunicação no enlace de descida (*downlink*), ilustrado na Fig. 44, em que uma estação rádio base (*Base Station* (BS)) equipada com N_t antenas atende simultaneamente K usuários com antena única. O conjunto total de usuários é definido como

$$\mathcal{K} = \{1, 2, \dots, K\}.$$

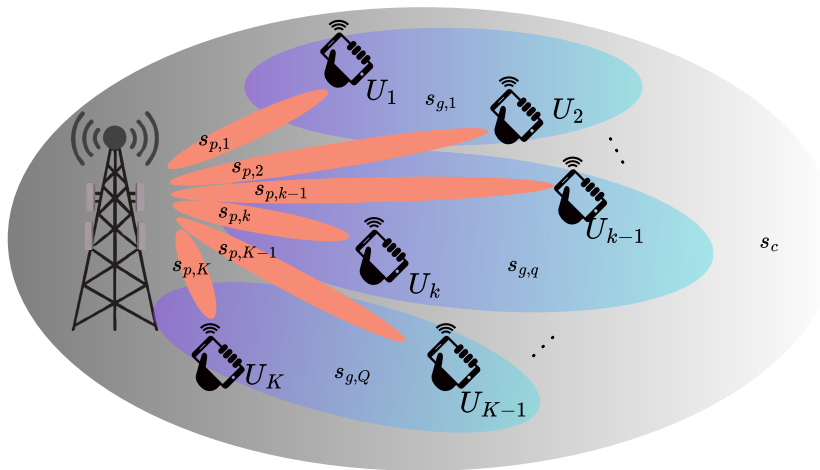


Figura 44: Modelo do sistema.

Os usuários são organizados em Q grupos distintos, indexados por $q \in \mathcal{Q} = \{1, 2, \dots, Q\}$. Cada grupo contém um subconjunto de usuários $\mathcal{K}_q$, de forma que os grupos sejam mutuamente exclusivos, isto é,

$$\mathcal{K}_q \cap \mathcal{K}_i = \emptyset, \quad q \neq i,$$

e a união de todos os grupos corresponda ao conjunto completo de usuários:

$$\bigcup_{q \in \mathcal{Q}} \mathcal{K}_q = \mathcal{K}.$$

O sistema emprega um esquema de HRSMA, no qual as mensagens dos usuários são divididas em múltiplas camadas de informação utilizando SIC [157]. Diferentemente do RSMA convencional, que utiliza apenas uma camada comum e uma camada privada, o HRSMA introduz uma estrutura hierárquica composta pelas seguintes mensagens:

- Uma mensagem comum global, destinada a todos os usuários;
- Mensagens comuns de grupo, destinadas apenas aos usuários pertencentes ao mesmo grupo;
- Mensagens privadas individuais.

Essa arquitetura é especialmente adequada para cenários com usuários agrupados geograficamente ou funcionalmente, como aplicações *Vehicle-to-Everything* (V2X), redes industriais ou sensores *Internet of Things* (IoT), permitindo maior eficiência espectral e melhor gerenciamento da interferência [149].

O processo de codificação ocorre da seguinte forma. A mensagem W_k , destinada ao usuário $k \in \mathcal{K}_q$, é dividida em três partes independentes:

$$W_k \rightarrow \{W_{c,k}, W_{g,q,k}, W_{p,k}\},$$

em que

- $W_{c,k}$ corresponde à parte comum global;
- $W_{g,q,k}$ representa a parte comum do grupo;
- $W_{p,k}$ corresponde à parte privada.

As partes comuns globais de todos os usuários são agregadas em uma única mensagem comum global W_c , posteriormente codificada no fluxo s_c , que será decodificado por todos os usuários do sistema.

Da mesma forma, para cada grupo q , as mensagens comuns de grupo são combinadas em uma única mensagem $W_{g,q}$, codificada no fluxo $s_{g,q}$, destinada exclusivamente aos usuários pertencentes ao grupo q .

Por fim, cada mensagem privada $W_{p,k}$ é codificada individualmente no fluxo $s_{p,k}$, destinado exclusivamente ao usuário k .

Assim, o vetor de sinal transmitido pela BS, denotado por $\mathbf{x} \in \mathbb{C}^{N_t}$, é formado pela superposição linear de todos os fluxos transmitidos:

$$\mathbf{x} = \sqrt{P\alpha_c} \mathbf{w}_c s_c + \sum_{i \in \mathcal{Q}} \sqrt{P\alpha_{g,i}} \mathbf{w}_{g,i} s_{g,i} + \sum_{j \in \mathcal{K}} \sqrt{P\alpha_{p,j}} \mathbf{w}_{p,j} s_{p,j}, \quad (165)$$

em que

- P representa a potência total de transmissão;
- α_c , $\alpha_{g,i}$ e $\alpha_{p,j}$ representam os coeficientes de alocação de potência associados às mensagens comum global, comuns de grupo e privadas, respectivamente;
- $\mathbf{w}_c$, $\mathbf{w}_{g,i}$ e $\mathbf{w}_{p,j}$ representam os vetores de pré-codificação associados aos respectivos fluxos;
- s_c , $s_{g,i}$ e $s_{p,j}$ representam os símbolos transmitidos.

Os coeficientes de potência devem satisfazer a restrição

$$\alpha_c + \sum_{i \in \mathcal{Q}} \alpha_{g,i} + \sum_{j \in \mathcal{K}} \alpha_{p,j} \leq 1.$$

O sinal recebido pelo usuário k pertencente ao grupo q é dado por

$$y_k = \mathbf{h}_k^H \mathbf{x} + n_k, \quad (166)$$

em que

- $\mathbf{h}_k \in \mathbb{C}^{N_t}$ representa o vetor de canal entre a BS e o usuário k ;
- $(\cdot)^H$ representa a operação hermitiana;
- $n_k \sim \mathcal{CN}(0, N_0)$ corresponde ao *Additive White Gaussian Noise* (AWGN).

Assume-se que o canal sofre desvanecimento Nakagami- m , amplamente utilizado para modelar diferentes condições de propagação sem fio. O ganho de potência do canal é definido como

$$g_k = \|\mathbf{h}_k\|^2,$$

e segue uma distribuição Gamma:

$$g_k \sim \Gamma(A_k, B_k),$$

em que

$$A_k = N_t m_k,$$

corresponde ao parâmetro de forma da distribuição, enquanto

$$B_k = \frac{\Omega_k}{m_k},$$

representa o parâmetro de escala, considerando componentes de desvanecimento independentes e identicamente distribuídas, obtido via casamento de momentos [151].

O parâmetro $m_k \geq 0.5$ representa a severidade do desvanecimento Nakagami e Ω_k corresponde à potência média das componentes multipercurso.

Expandindo a Eq. (165), o sinal recebido pode ser reescrito explicitamente destacando as componentes desejadas e as interferências:

$$\begin{aligned}
 y_k = & \underbrace{\sqrt{P\alpha_c} \mathbf{h}_k^H \mathbf{w}_c s_c}_{\text{Mensagem comum global}} + \underbrace{n_k}_{\text{Ruído}} \\
 & + \underbrace{\sqrt{P\alpha_{g,q}} \mathbf{h}_k^H \mathbf{w}_{g,q} s_{g,q}}_{\text{Mensagem comum do grupo}} \\
 & + \underbrace{\sum_{i \in \mathcal{Q}, i \neq q} \sqrt{P\alpha_{g,i}} \mathbf{h}_k^H \mathbf{w}_{g,i} s_{g,i}}_{\text{Interferência entre grupos}} \\
 & + \underbrace{\sqrt{P\alpha_{p,k}} \mathbf{h}_k^H \mathbf{w}_{p,k} s_{p,k}}_{\text{Mensagem privada}} \\
 & + \underbrace{\sum_{j \in \mathcal{K}, j \neq k} \sqrt{P\alpha_{p,j}} \mathbf{h}_k^H \mathbf{w}_{p,j} s_{p,j}}_{\text{Interferência multiusuário}}. \tag{167}
 \end{aligned}$$

15.2.1 Decodificação dos Sinais e Análise da SINR

O processo de decodificação no receptor segue o princípio de SIC, no qual as mensagens são decodificadas sequencialmente, removendo-se as componentes já detectadas antes da próxima etapa de detecção.

Define-se inicialmente a relação sinal-ruído média de transmissão como

$$\rho = \frac{P}{N_0}.$$

Na primeira etapa, todos os usuários decodificam a mensagem comum global s_c . A *Signal to Interference plus Noise Ratio* (SINR) correspondente para o usuário k é dada por

$$\gamma_{c,k} = \frac{\rho \alpha_c |\mathbf{h}_k^H \mathbf{w}_c|^2}{\sum_{i \in \mathcal{Q}} \rho \alpha_{g,i} |\mathbf{h}_k^H \mathbf{w}_{g,i}|^2 + \sum_{j \in \mathcal{K}} \rho \alpha_{p,j} |\mathbf{h}_k^H \mathbf{w}_{p,j}|^2 + 1}. \tag{168}$$

A taxa instantânea associada à mensagem comum global é

$$R_{c,k} = \log_2(1 + \gamma_{c,k}).$$

Como todos os usuários precisam decodificar corretamente essa mensagem, a taxa efetiva do fluxo comum global é limitada pelo usuário com pior condição de canal:

$$R_c = \min_{k \in \mathcal{K}} \{R_{c,k}\}. \tag{169}$$

Após cancelar o fluxo comum global utilizando SIC, o usuário $k \in \mathcal{K}_q$ decodifica a mensagem comum do grupo $s_{g,q}$. A SINR correspondente é dada por

$$\gamma_{g,q,k} = \frac{\rho \alpha_{g,q} |\mathbf{h}_k^H \mathbf{w}_{g,q}|^2}{\sum_{i \in \mathcal{Q}, i \neq q} \rho \alpha_{g,i} |\mathbf{h}_k^H \mathbf{w}_{g,i}|^2 + \sum_{j \in \mathcal{K}} \rho \alpha_{p,j} |\mathbf{h}_k^H \mathbf{w}_{p,j}|^2 + 1}. \tag{170}$$

A taxa correspondente é expressa por

$$R_{g,q,k} = \log_2(1 + \gamma_{g,q,k}),$$

e a taxa efetiva do grupo é limitada pelo pior usuário pertencente ao grupo:

$$R_{g,q} = \min_{k \in \mathcal{K}_q} \{R_{g,q,k}\}. \quad (171)$$

Após remover as camadas comuns, o usuário decodifica sua mensagem privada $s_{p,k}$, tratando os fluxos remanescentes como interferência. A SINR correspondente é dada por

$$\gamma_{p,k} = \frac{\rho \alpha_{p,k} |\mathbf{h}_k^H \mathbf{w}_{p,k}|^2}{\sum_{i \in \mathcal{Q}, i \neq q} \rho \alpha_{g,i} |\mathbf{h}_k^H \mathbf{w}_{g,i}|^2 + \sum_{j \in \mathcal{K}, j \neq k} \rho \alpha_{p,j} |\mathbf{h}_k^H \mathbf{w}_{p,j}|^2 + 1}. \quad (172)$$

A taxa privada é definida como

$$R_{p,k} = \log_2(1 + \gamma_{p,k}).$$

Finalmente, a taxa total alcançável pelo usuário k é dada pela soma da taxa privada com as parcelas atribuídas das taxas comum global e comum de grupo:

$$R_k^{\text{tot}} = C_{c,k} + C_{g,k} + R_{p,k}, \quad (173)$$

em que $C_{c,k}$ e $C_{g,k}$ representam as parcelas das taxas comum global e comum de grupo atribuídas ao usuário k , respectivamente.

15.3 Análise da Probabilidade de Indisponibilidade

Nesta seção, são obtidas expressões em forma fechada para a OP, bem como uma análise assintótica no regime de alta *Signal-to-Noise Ratio* (SNR), a qual fundamenta a formulação de alocação de potência via PG.

A indisponibilidade ocorre quando qualquer uma das etapas do processo de SIC não atende à taxa mínima requerida R_i^{th} . Definindo os limiares de SINR como $\gamma_i^b = 2^{R_i^{\text{th}}} - 1$, para $i \in \{c, q, p\}$, a probabilidade de indisponibilidade do usuário $k \in \mathcal{K}$ é dada por

$$P_{\text{out},k} = 1 - \mathbb{P}(\gamma_{c,k} > \gamma_c^b, \gamma_{g,q,k} > \gamma_q^b, \gamma_{p,k} > \gamma_p^b). \quad (174)$$

Substituindo as expressões de SINR, obtêm-se as condições mínimas para o ganho de canal $g_k = \|\mathbf{h}_k\|^2$:

$$g_k > \frac{\gamma_c^b}{\rho [\alpha_c - \gamma_c^b (\alpha_{g,q} + \alpha_{p,k})]} \triangleq \Theta_{c,k} \quad (175)$$

$$g_k > \frac{\gamma_q^b}{\rho (\alpha_{g,q} - \gamma_q^b \alpha_{p,k})} \triangleq \Theta_{q,k} \quad (176)$$

$$g_k > \frac{\gamma_p^b}{\rho \alpha_{p,k}} \triangleq \Theta_{p,k}. \quad (177)$$

A condição conjunta de sucesso é dada por

$$g_k > \Theta_k = \max\{\Theta_{c,k}, \Theta_{q,k}, \Theta_{p,k}\}. \quad (178)$$

Logo, a OP pode ser expressa como

$$P_{out,k} = F_{g_k}(\Theta_k). \quad (179)$$

Considerando distribuição Gamma,

$$P_{out,k} = 1 - \frac{\Gamma(A_k, \Theta_k/B_k)}{\Gamma(A_k)}. \quad (180)$$

No regime de alta SNR, obtém-se

$$\bar{P}_{out,k} \approx \frac{1}{\Gamma(A_k + 1)} \left(\frac{\Theta_k}{B_k} \right)^{A_k}. \quad (181)$$

Separando o efeito da SNR, resulta que

$$\bar{P}_{out,k} \approx \frac{1}{\Gamma(A_k + 1)} \left(\frac{\varrho_k}{B_k \rho} \right)^{A_k}, \quad (182)$$

em que

- Ganho de diversidade: $G_d = A_k = N_t m_k$
- Ganho de codificação: $G_c = \frac{B_k}{\varrho_k}$

Observa-se que o ganho de diversidade depende apenas do canal e do número de antenas, enquanto o ganho de codificação é influenciado pela alocação de potência.

15.4 Alocação de Potência Min–Max Baseada em Programação Geométrica

A análise assintótica apresentada na Seção 15.3 conduz a uma função objetivo tratável, uma vez que a estrutura monomial da expressão da OP permite sua otimização por meio de PG. Com o objetivo de garantir equidade ao longo da cadeia hierárquica de decodificação, propõe-se uma estratégia de alocação de potência do tipo min–max, que minimiza a pior probabilidade de indisponibilidade entre todos os usuários.

15.4.1 Formulação do Problema

O problema de otimização consiste em minimizar a maior OP entre os usuários $k \in \mathcal{K}$, por meio do ajuste do vetor de alocação de potência α :

$$P_1 : \min_{\alpha} \left(\max_{k \in \mathcal{K}} \bar{P}_{out,k}(\alpha) \right), \quad (183)$$

sujeito às seguintes restrições:

$$\alpha_c + \sum_{q \in \mathcal{Q}} \alpha_{g,q} + \sum_{k \in \mathcal{K}} \alpha_{p,k} \leq 1, \quad (184)$$

$$\alpha_{g,q} > \gamma_{b,q} \alpha_{p,k}, \quad \forall k \in \mathcal{K}_q, \quad (185)$$

$$\alpha_c > \gamma_{b,c} (\alpha_{g,q} + \alpha_{p,k}), \quad \forall k \in \mathcal{K}_q, \quad (186)$$

$$\alpha_c, \alpha_{g,q}, \alpha_{p,k} > 0. \quad (187)$$

A restrição de soma assegura o limite total de potência disponível, enquanto as demais garantem a viabilidade do processo de SIC, impondo margens positivas entre os níveis de potência das diferentes camadas de transmissão.

15.4.2 Reformulação via Programação Geométrica

Para tornar o problema de otimização tratável no contexto de PG, introduz-se uma variável auxiliar t (epígrafo), responsável por limitar superiormente a OP de todos os usuários. Dessa forma, o problema min-max pode ser reescrito como

$$P_{PG} : \min_{\alpha} t \quad (188)$$

sujeito a

$$\bar{P}_{\text{out},k}(\alpha) \leq t, \quad \forall k \in \mathcal{K}, \quad (189)$$

$$\alpha_c + \sum_{q \in \mathcal{Q}} \alpha_{g,q} + \sum_{k \in \mathcal{K}} \alpha_{p,k} \leq 1. \quad (190)$$

A partir da expressão assintótica da OP, observa-se que o limiar dominante Θ_k pode ser decomposto em três contribuições associadas às etapas de decodificação: (i) camada comum do sistema, (ii) camada comum de grupo e (iii) camada privada. Essa decomposição permite reescrever as restrições de desempenho em termos de funções posinomiais.

Definindo constantes normalizadas $K_{i,k}$, para $i \in \{c, q, p\}$, que agregam os parâmetros do canal, os requisitos de taxa e o regime de SNR, obtêm-se as seguintes restrições equivalentes para cada usuário $k \in \mathcal{K}_q$:

$$\gamma_{b,c}(\alpha_{g,q} + \alpha_{p,k})\alpha_c^{-1} + K_{c,k} t^{-1/A_k} \alpha_c^{-1} \leq 1, \quad (191)$$

$$\gamma_{b,q} \alpha_{p,k} \alpha_{g,q}^{-1} + K_{g,k} t^{-1/A_k} \alpha_{g,q}^{-1} \leq 1, \quad (192)$$

$$K_{p,k} t^{-1/A_k} \alpha_{p,k}^{-1} \leq 1. \quad (193)$$

As restrições acima possuem forma posinomial, enquanto a função objetivo é monomial. Consequentemente, o problema resultante configura um PG padrão, o que possibilita sua resolução eficiente por métodos convexos.

A solução ótima pode ser obtida por meio do *framework CVX*, em conjunto com o solucionador MOSEK. Nesse processo, o CVX realiza automaticamente as transformações logarítmicas necessárias para converter o problema em sua forma convexa equivalente, enquanto o MOSEK resolve o problema resultante com elevada eficiência computacional.

A solução α^* obtida fornece a alocação de potência que minimiza a pior OP entre os usuários, garantindo simultaneamente desempenho robusto e equidade no sistema HRSMA.

15.5 Resultados Numéricos

Nesta seção, avaliamos o desempenho do esquema HRSMA utilizando as expressões analíticas obtidas, as quais são validadas por meio de simulações de Monte Carlo. Além disso, comparamos a estratégia de alocação de potência proposta baseada em PG com um cenário de referência, no qual são adotados coeficientes arbitrários dados por $\alpha_c = 1/3$, $\alpha_{g,q} = 1/(3Q)$ e $\alpha_{p,k} = 1/(3K)$.

Salvo indicação em contrário, os resultados numéricos consideram um canal com desvanecimento Nakagami- m , com parâmetro $m_k = 1.5$, potência média dos múltiplos percursos $\Omega_k = 1$, e taxas alvo $\{R_c, R_{g,q}, R_{p,k}\} = \{0.5, 0.8, 1.0\}$ bps/Hz, para todo $k \in \mathcal{K}$ e $q \in \mathcal{Q}$. Observa-se uma excelente concordância entre as curvas analíticas e os resultados de Monte Carlo ao longo de toda a faixa de SNR, o que confirma a validade das expressões em forma fechada derivadas para o modelo HRSMA.

Para facilitar a visualização, as figuras a seguir apresentam apenas a probabilidade de indisponibilidade do usuário com pior desempenho. Em cenários com condições de canal heterogêneas ou requisitos de taxa assimétricos, os usuários naturalmente apresentam diferentes desempenhos de *outage*, o que normalmente resultaria em múltiplas curvas. No entanto, os resultados demonstram que a estratégia de otimização min-max baseada em PG promove equidade entre os usuários. Ao priorizar o pior caso, o algoritmo ajusta a alocação de potência de forma que as probabilidades de *outage* de todos os K usuários converjam para um mesmo valor. Como consequência, as curvas otimizadas se sobrepõem perfeitamente, sendo representadas por uma única curva para comparação com o caso de alocação arbitrária.

A Fig. 45 ilustra a probabilidade de *outage* em cenários homogêneos e heterogêneos, considerando $N_t = 4$, $K = 4$ e $Q = 2$. No cenário homogêneo, assume-se $\Omega_k = 1$ para todos os usuários, enquanto no cenário heterogêneo consideram-se dois usuários com canais fortes ($\Omega_k = 2$) e dois com canais mais fracos ($\Omega_k = 1$). São avaliadas duas estratégias de agrupamento: (i) agrupamento homogêneo, que reúne usuários com condições de canal semelhantes, e (ii) agrupamento heterogêneo, que combina usuários fortes e fracos.

Observa-se que, no caso heterogêneo, o agrupamento homogêneo apresenta desempenho superior, pois permite uma distribuição de potência mais equilibrada e evita que usuários mais fracos dominem o desempenho do grupo. Já no cenário com alocação arbitrária, as curvas se sobrepõem, uma vez que o desempenho é limitado pelo usuário mais fraco, independentemente da estratégia de agrupamento. As curvas otimizadas via PG apresentam ganhos de codificação significativos, mantendo a mesma ordem de diversidade, como evidenciado pelas inclinações assintóticas paralelas.

A Fig. 46 analisa o impacto do número de antenas, comparando os casos $N_t = 4$ e $N_t = 8$, com $K = 4$ e $Q = 2$. Observa-se que o aumento no número de antenas resulta em uma inclinação mais acentuada das curvas, indicando um maior ganho de diversidade espacial. Além disso, a otimização baseada em PG proporciona um ganho adicional em SNR, evidenciando que a combinação de diversidade espacial e alocação eficiente de potência melhora significativamente a confiabilidade do sistema.

Na Fig. 47, comparamos duas configurações de agrupamento para um sistema maior ($N_t = 8$, $K = 8$): (i) vários grupos pequenos ($Q = 4$, $K_q = 2$) e (ii) poucos grupos maiores ($Q = 2$, $K_q = 4$). Embora a ordem de diversidade permaneça inalterada, observa-se que grupos maiores apresentam melhor desempenho, sugerindo uma mitigação mais eficiente da interferência via camada *group-common*. A otimização via PG novamente proporciona ganhos adicionais de SNR.

Por fim, a Fig. 48 avalia o impacto do número total de usuários ($K = 4, 6, 8$), considerando $N_t = 8$ e $Q = 2$, resultando em tamanhos de grupo $K_q = 2, 3, 4$. Nota-se que o desempenho se degrada à medida que o número de usuários por grupo aumenta, devido ao aumento da interferência intra-grupo e da competição por recursos. Entretanto, a inclinação assintótica das curvas permanece constante, indicando que a ordem de diversidade não é afetada. Além disso, a otimização baseada em PG demonstra robustez mesmo em cenários com maior densidade de usuários, sendo capaz de superar configurações arbitrárias com menor número de usuários.

De forma geral, os resultados evidenciam que a escolha adequada da estratégia de agrupamento e a utilização de alocação de potência otimizada são fundamentais para maximizar a eficiência e a confiabilidade do sistema HRSMA.

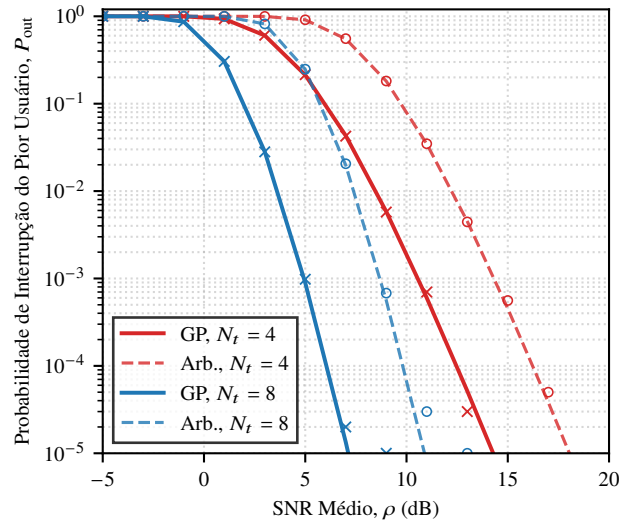


Figura 45: Probabilidade de indisponibilidade do pior usuário em função da SNR de transmissão para cenários homogêneos e heterogêneos com $N_t = 4$, $K = 4$ e $Q = 2$, comparando a alocação de potência otimizada via PG com uma abordagem arbitrária.

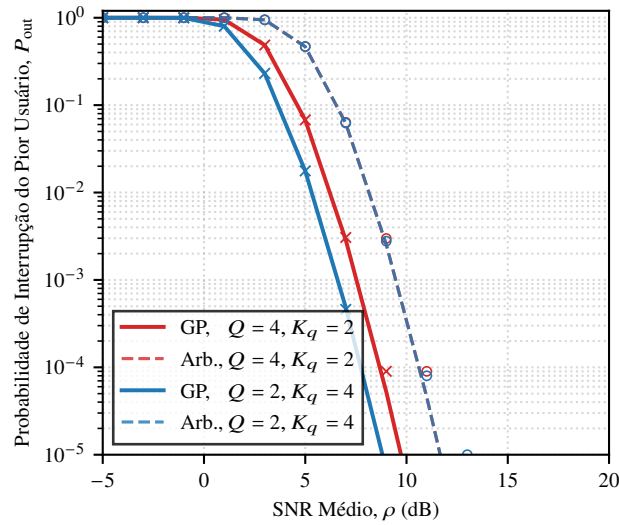


Figura 46: Probabilidade de indisponibilidade do pior usuário para $N_t = 4$ e $N_t = 8$ antenas, com $K = 4$ e $Q = 2$, evidenciando o ganho de diversidade espacial e o ganho adicional proporcionado pela otimização via PG.

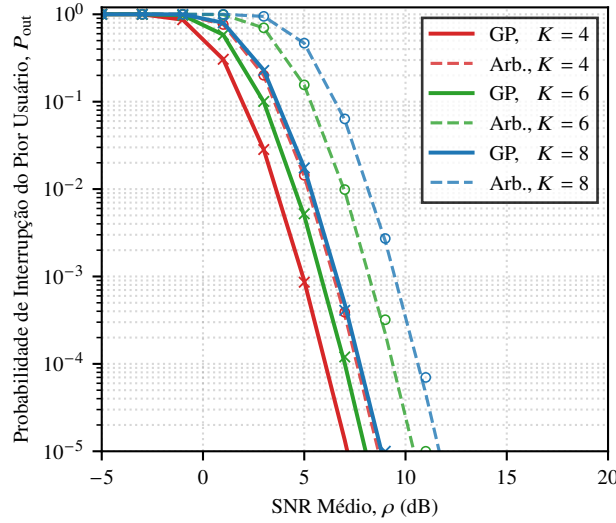


Figura 47: Probabilidade de indisponibilidade do pior usuário para diferentes configurações de agrupamento: 4 grupos com 2 usuários ($Q = 4$, $K_q = 2$) versus 2 grupos com 4 usuários ($Q = 2$, $K_q = 4$), com $N_t = 8$ e $K = 8$.

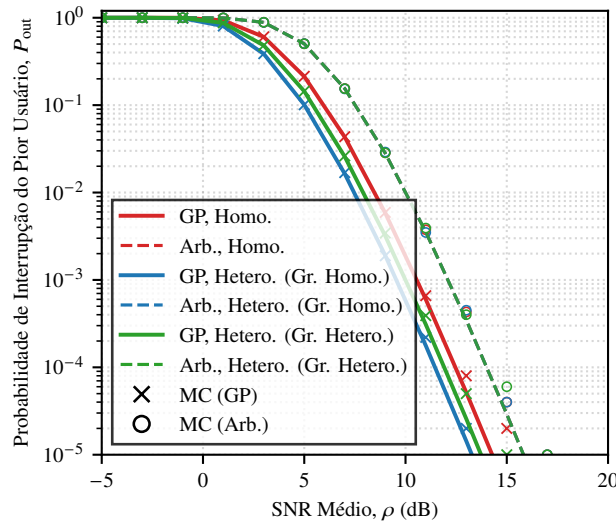


Figura 48: Probabilidade de indisponibilidade do pior usuário para diferentes números de usuários $K = 4, 6$ e 8 , com $N_t = 8$ e $Q = 2$, resultando em tamanhos de grupo $K_q = 2, 3$ e 4 , respectivamente.

15.6 Conclusão

Este trabalho analisou o desempenho do HRSMA em enlaces *downlink* sob desvanecimento Nakagami- m . Utilizando um esquema de pré-codificação híbrido ZF-MRT, foram derivadas expressões exatas em forma fechada para a (OP), bem como sua caracterização assintótica em regime de alta SNR. Os resultados evidenciam que o ganho de diversidade é determinado pela configuração do arranjo de antenas e pelas condições de desvanecimento, enquanto o ganho de codificação é diretamente influenciado pela estratégia de alocação de potência.

Adicionalmente, foi proposta uma abordagem de alocação de potência do tipo min-max baseada em PG, com o objetivo de garantir equidade entre os usuários. Os resultados numéricos corroboram a análise teórica, demonstrando ganhos expressivos de codificação e fornecendo importantes diretrizes de projeto. Em particular, observa-se que o agrupamento homogêneo supera o agrupamento heterogêneo em termos de desempenho, enquanto configurações com menor número de grupos e maior quantidade de usuários por grupo tendem a oferecer maior confiabilidade.

Por outro lado, o aumento do tamanho dos grupos implica degradação de desempenho devido ao aumento da interferência intra-grupo. Ainda assim, a otimização baseada em PG é capaz de mitigar significativamente esse efeito, evidenciando a robustez do HRSMA mesmo em cenários com alta densidade de usuários.

16 Desempenho de Sistemas NTN Baseados em NOMA sob Desvanecimento Generalizado

Francisco Raimundo Albuquerque Parente, Bruna do Nascimento Benevides, Raphael Parreira de Souza e Maykon Renan Pereira da Silva
CPQD

16.1 Introdução

As redes *Non-Terrestrial Networks* (NTN) despontam como um dos pilares dos futuros sistemas Sexta Geração de Rede Móvel Celular (6G) ao expandirem a conectividade para além da infraestrutura terrestre, viabilizando cobertura em áreas remotas, cenários de desastre e ambientes com baixa densidade de estações rádio-base [158, 159]. Ao integrarem satélites, em especial constelações *Low Earth Orbit* (LEO), plataformas de alta altitude e *Unmanned Aerial Vehicle* (UAV), as redes NTN passam a operar sob condições de enlace significativamente distintas daquelas observadas em redes convencionais, incluindo maiores variações de distância, elevada mobilidade relativa, efeitos de múltiplos feixes e regimes de propagação caracterizados por componentes *Line of Sight* (LOS) dominantes combinadas com multipercursos e sombreamento [160, 161, 162]. Essa heterogeneidade torna a análise de confiabilidade e eficiência espectral particularmente desafiadora, uma vez que a experiência do usuário passa a ser fortemente influenciada por flutuações rápidas do canal e por assimetrias entre enlaces, características típicas de cenários NTN [163].

Nesse contexto, as redes *Non-Orthogonal Multiple Access* (NOMA) têm sido amplamente investigadas como uma tecnologia habilitadora para NTN, por viabilizarem a multiplexação simultânea de múltiplos usuários em um mesmo recurso de rádio no domínio tempo-frequência, explorando o domínio da potência e a decodificação por meio de *Successive Interference Cancellation* (SIC) [164, 165]. Em cenários NTN, o NOMA se destaca principalmente por dois potenciais benefícios: (i) elevar a eficiência espectral em ambientes com recursos limitados e (ii) ampliar a conectividade para usuários com condições de enlace heterogêneas, explorando diferenças de ganho de canal [166]. Infelizmente, o desempenho do NOMA é altamente sensível ao comportamento estatístico do canal, incluindo correlação, componentes especulares, número de *clusters* e não-linearidades, bem como às imperfeições do SIC, associadas à interferência residual. Por essa razão, tornam-se necessários modelos de canal capazes de representar, de forma realista, a diversidade de condições encontradas em ambientes NTN. Assim, análises fundamentadas em modelos excessivamente restritos podem conduzir a conclusões demasiadamente otimistas ou com limitada capacidade de generalização, o que reforça a necessidade do uso de distribuições mais abrangentes para a caracterização da Probabilidade de Indisponibilidade (*Outage Probability*), da capacidade e da robustez do sistema [167].

16.2 Modelo de Canal de Desvanecimento

Em ambientes caracterizados por múltiplos percursos de propagação, o sinal recebido pode ser modelado como a superposição de diversas componentes com amplitudes e fases aleatórias. De acordo com o *Central Limit Theorem* (CLT), a soma de um número suficientemente grande de componentes independentes tende a apresentar comportamento Gaussiano. Esse princípio fundamenta os modelos de desvanecimento da classe Gaussiana, nos quais os diferentes *clusters* de propagação são representados como combinações de variáveis aleatórias Gaussianas.

Neste trabalho, adota-se um modelo geral de desvanecimento capaz de capturar múltiplos *clusters*, componentes especulares e efeitos de não linearidade, sendo adequado para cenários NTN-NOMA. A potência do canal B é inicialmente definida como a soma de M variáveis aleatórias Gaussianas quadráticas, dada por

$$B^{\frac{\alpha}{2}} = \sum_{i=1}^M X_i^2, \quad (194)$$

em que $\alpha > 0$ representa o parâmetro de não linearidade do canal, e M corresponde ao número de componentes Gaussianas. O vetor aleatório $\mathbf{X} \triangleq [X_1, X_2, \dots, X_M]^T$ segue uma *Probability Density Function* (PDF) Gaussiana multivariada com vetor média $\mathbf{m} \triangleq \mathbb{E}[\mathbf{X}]$ e matriz de covariância $\mathbf{\Sigma}$ [168].

Os elementos do vetor média são dados por $\mathbb{E}[X_i] = m_i$, enquanto as variâncias são $\mathbb{V}[X_i] = \sigma_i^2$, $\forall i \in \mathcal{M}$, com $\mathcal{M} \triangleq \{1, \dots, M\}$. A matriz de covariância, assumida positiva definida, pode ser expressa como

$$\mathbf{\Sigma} = \begin{bmatrix} \sigma_1^2 & \rho_{1,2}\sigma_1\sigma_2 & \cdots & \rho_{1,M}\sigma_1\sigma_M \\ \rho_{1,2}\sigma_1\sigma_2 & \sigma_2^2 & \cdots & \rho_{2,M}\sigma_2\sigma_M \\ \vdots & \vdots & \ddots & \vdots \\ \rho_{1,M}\sigma_1\sigma_M & \rho_{2,M}\sigma_2\sigma_M & \cdots & \sigma_M^2 \end{bmatrix}, \quad (195)$$

em que $\rho_{i,j} \equiv \rho(X_i, X_j)$ representa o coeficiente de correlação entre as variáveis X_i e X_j .

A PDF conjunta de $\mathbf{X}$ é então dada por

$$f_{\mathbf{X}}(\mathbf{x}) = \frac{\exp\left[-\frac{1}{2}(\mathbf{x} - \mathbf{m})^T \mathbf{\Sigma}^{-1}(\mathbf{x} - \mathbf{m})\right]}{[(2\pi)^M \det(\mathbf{\Sigma})]^{\frac{1}{2}}}, \quad (196)$$

em que $\mathbf{x} \triangleq [x_1, x_2, \dots, x_M]^T$.

A análise exata do desempenho do sistema, em termos de Probabilidade de Indisponibilidade, requer a avaliação de integrais multidimensionais sobre a PDF Gaussiana multivariada apresentada em (196). Nesse contexto, a PDF da potência do canal associada ao modelo em (194) pode ser expressa como

$$f_B(\beta) = \alpha \beta^{\frac{\alpha}{2}-1} \left[(2^{\frac{2}{M}+3} \pi)^M \det(\mathbf{\Sigma}) \right]^{-\frac{1}{2}} \quad (197)$$

$$\times \int_0^{\beta^{\frac{\alpha}{2}}} \int_0^{\beta^{\frac{\alpha}{2}} - s_M} \cdots \int_0^{\beta^{\frac{\alpha}{2}} - \sum_{i=3}^M s_i} \psi(\mathbf{s}) ds_2 \cdots ds_M, \quad (198)$$

em que

$$\psi(\mathbf{s}) \triangleq \left(\prod_{i=1}^M \sqrt{s_i} \right)^{-1} \sum_{j=1}^{2^M} \exp \left[-\frac{1}{2}(\mathbf{r}_j - \mathbf{m})^T \mathbf{\Sigma}^{-1}(\mathbf{r}_j - \mathbf{m}) \right], \quad (199)$$

com $s_i \triangleq x_i^2$, $\mathbf{s} \triangleq [s_1, s_2, \dots, s_M]^T$ e $s_1 = \beta^{\alpha/2} - \sum_{i=2}^M s_i$. O vetor $\mathbf{r}_j$ representa a j -ésima combinação possível dos sinais das componentes $\{\sqrt{s_i}\}$, totalizando 2^M combinações [168].

A métrica de desempenho considerada é a Probabilidade de Indisponibilidade (P_{out}), definida como a probabilidade de que a *Signal-to-Noise Ratio* (SNR) instantânea γ seja inferior a um limiar γ_{th} , isto é,

$$P_{\text{out}} \triangleq P[\gamma < \gamma_{\text{th}}] = \int_0^{\frac{\gamma_{\text{th}}}{\alpha}} f_B(\beta) d\beta, \quad (200)$$

em que $\bar{\gamma}$ representa a SNR média, assumindo $\mathbb{E}[B] = 1$.

Adicionalmente, essa métrica pode ser analisada sob a perspectiva dos ganhos de diversidade e de codificação, os quais são diretamente influenciados pelos parâmetros do modelo de desvanecimento considerado. Esses aspectos serão explorados nas seções subsequentes.

16.3 Modelo do Sistema

O modelo do sistema e a formulação da *Signal to Interference plus Noise Ratio* (SINR) adotados neste trabalho, incluindo os princípios do NOMA, a estratégia de alocação de potência e a consideração de imperfeições de *hardware* e de cancelamento de interferência, são baseados em [169].

Em contraste com os esquemas de acesso múltiplo ortogonal, nos quais os recursos de tempo e/ou frequência são alocados de forma exclusiva entre os usuários, o NOMA explora o compartilhamento simultâneo desses recursos. Nessa abordagem, a distinção entre os sinais é realizada no domínio da potência, permitindo maior aproveitamento espectral e suportando um número elevado de dispositivos conectados. No cenário de transmissão no enlace *downlink*, o funcionamento do NOMA baseia-se na aplicação de *Superposition Coding* (SC) no transmissor, enquanto os receptores empregam SIC para a separação sucessiva dos sinais superpostos.

16.3.1 *Superposition Coding*

No transmissor de um sistema NOMA, os dados destinados a múltiplos usuários são agregados em um único sinal transmitido pela interface aérea. Para que o receptor consiga recuperar corretamente cada informação, a unidade de processamento em banda base distribui a potência de forma distinta entre os sinais. Considerando que o transmissor disponha de uma potência total P e que p_1 e p_2 representem os coeficientes de alocação de potência dos Usuários 1 e 2, respectivamente, o sinal transmitido para o caso de dois usuários pode ser descrito como

$$x = \sqrt{p_1 P} s_1 + \sqrt{p_2 P} s_2. \quad (201)$$

Essa estratégia de alocação de potência é essencial para o desempenho do sistema e não ocorre de forma arbitrária. De modo geral, usuários com condições de canal mais desfavoráveis, tipicamente mais distantes do transmissor, recebem uma maior fração da potência disponível. Por outro lado, usuários com melhores condições de canal, geralmente mais próximos, são atendidos com uma menor parcela de potência.

16.3.2 *Successive Interference Cancellation*

No receptor, o sinal transmitido x é observado de maneira distinta por cada usuário, uma vez que ele é afetado pelas condições individuais de canal h_k e pela adição de *Additive White Gaussian Noise* (AWGN) n_k . Para um Usuário k , o sinal recebido y_k pode ser expresso como

$$y_k = h_k^H x + n_k = h_k^H (\sqrt{p_1 P} s_1 + \sqrt{p_2 P} s_2) + n_k. \quad (202)$$

De acordo com a estratégia de alocação de potência descrita na subseção anterior, o usuário mais distante recebe o sinal superposto e trata o sinal indesejado como interferência adicional em sua medida de SINR. Por outro lado, o usuário com melhor condição de canal recebe o sinal destinado ao usuário mais distante com qualidade suficiente para decodificá-lo primeiro e, em seguida, realizar o cancelamento sucessivo da interferência, de modo que reste apenas a sua própria mensagem.

16.3.3 Imperfeições do SIC e degradação de *hardware*

Em aplicações práticas, restrições impostas pelo hardware degradam a precisão das estimativas de canal, o que se reflete no desempenho da SINR dos usuários. Neste estudo, adota-se um sistema NOMA no *downlink* composto por dois usuários.

Define-se a SNR de transmissão como $\xi = \frac{P}{N_0}$, sendo P a potência total transmitida pelo UAV e N_0 a densidade espectral de potência do ruído. Nesse contexto, a SINR do usuário localizado em piores condições de canal pode ser representada por

$$\gamma_1 = \frac{\xi |h_1|^2 p_1}{\xi |h_1|^2 (p_2 + \kappa^2) + 1}, \quad (203)$$

em que $0 \leq \kappa \leq 1$ representa o nível de degradação de *hardware*. Quando $\kappa = 0$, assume-se a ausência de falhas de *hardware*; por outro lado, quando $\kappa = 1$, considera-se um cenário de degradação máxima.

Além disso, a hipótese de cancelamento perfeito é, em geral, irrealista em sistemas reais. Considerando o mesmo cenário, a SINR do usuário mais próximo pode ser expressa como

$$\gamma_2 = \frac{\xi |h_2|^2 p_2}{\xi |h_2|^2 (p_1 \epsilon + \kappa^2) + 1}, \quad (204)$$

em que $0 \leq \epsilon \leq 1$ denota o grau de imperfeição do SIC. Quando $\epsilon = 0$, assume-se cancelamento perfeito; quando $\epsilon = 1$, considera-se que não há cancelamento do sinal do Usuário 1.

16.3.4 Geometria do modelo e caracterização do canal

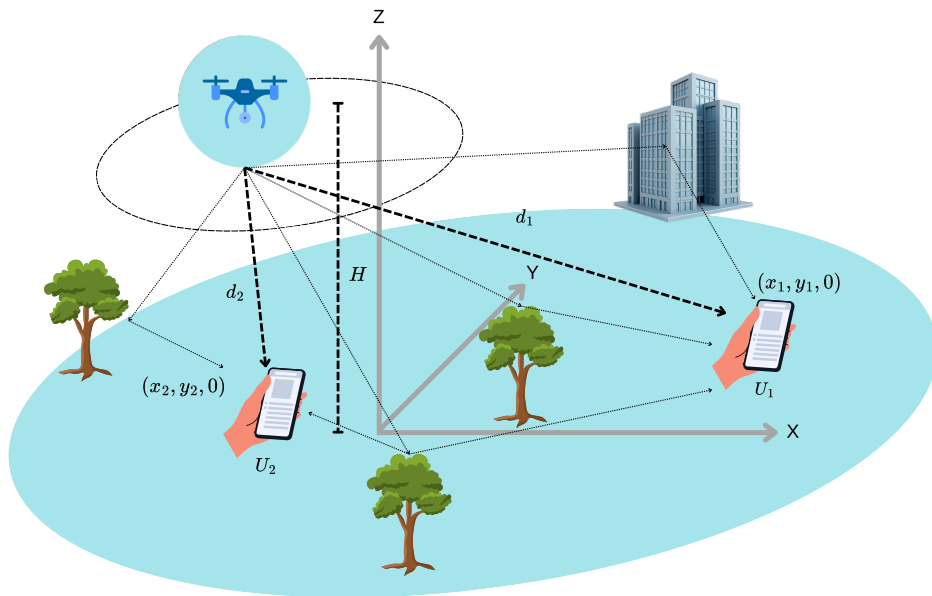


Figura 49: Modelo do sistema NTN baseado em NOMA.

No modelo proposto em [169], o UAV segue uma trajetória circular a uma altitude fixa H e com raio r_1 , fornecendo cobertura a usuários distribuídos aleatoriamente em uma região de

raio r_2 . A relação espacial entre o nó aéreo localizado em (x, y, z) e um usuário posicionado em $(x_m, y_m, 0)$ é descrita pela distância euclidiana

$$d_m = \sqrt{(x_m - x)^2 + (y_m - y)^2 + H^2}. \quad (205)$$

Para modelar a atenuação de larga escala, adota-se o modelo de *Free Space Path Loss* (FSPL), adequado a cenários com predominância de LOS, típicos de comunicações entre UAV e usuários terrestres em baixas altitudes, como ilustrado na Figura 49. A perda de percurso, em dB, é dada por

$$PL_{FSPL}(dB) = 32.4 + 20 \log_{10}(f_{MHz}) + 20 \log_{10}(d_{km}), \quad (206)$$

em que f_{MHz} representa a frequência de operação em MHz e d_{km} corresponde à distância d expressa em quilômetros. Para a simulação do ganho de canal, a perda de percurso é convertida para escala linear como

$$L(d) = 10^{-\frac{PL_{FSPL}(dB)}{10}}. \quad (207)$$

Define-se ainda um termo de normalização correspondente à perda em uma distância de referência de 1 km, dado por

$$L_0 = 10^{-\frac{PL_0(dB)}{10}}, \quad (208)$$

em que

$$PL_0(dB) = 32.4 + 20 \log_{10}(f_{MHz}). \quad (209)$$

O ganho de potência do canal para cada usuário é então modelado como

$$|h|^2 = \frac{L(d)}{L_0} B, \quad (210)$$

em que B representa o termo de desvanecimento de pequena escala, definido de acordo com o modelo de canal apresentado na Seção 16.2.

Por fim, dentre os usuários considerados, define-se como canal primário aquele associado ao menor ganho de canal e como canal secundário aquele associado ao maior ganho de canal, permitindo a avaliação do desempenho do sistema sob distintas condições de propagação.

16.4 Resultados Numéricos e Discussão

A Tabela 13 apresenta os parâmetros de simulação adotados nos experimentos numéricos. Esses parâmetros foram definidos de modo a representar, de forma realista, o ambiente de comunicação entre um UAV e usuários terrestres, considerando aspectos físicos do enlace e elementos práticos do sistema. O número de amostras de Monte Carlo foi fixado em 10.000, a fim de garantir a convergência estatística dos resultados e reduzir a variabilidade amostral associada à geração aleatória das posições dos usuários e das realizações do canal.

Os parâmetros geométricos do sistema, tais como o raio da trajetória do UAV (2,0 km), o raio da região de usuários (600 m) e a altura de voo (600 m), definem a topologia espacial da simulação e influenciam diretamente a distribuição das distâncias entre o UAV e os usuários. As taxas-alvo dos usuários primários e secundários, ambas fixadas em 0,5, representam os requisitos mínimos de desempenho considerados para cada enlace no cenário avaliado. A faixa de SNR

Tabela 13: Parâmetros de simulação do sistema NTN baseado em NOMA.

Parâmetro de simulação	Valor	Parâmetro de simulação	Valor
Amostras de Monte Carlo	10000	Taxa-alvo do usuário secundário	0,5
Raio do UAV	2,0	Degradações residuais de <i>hardware</i>	0,05
Raio da região dos usuários	600	SIC imperfeita residual	0,05
Altura do UAV	600	Coefficiente de potência do usuário primário	0,8
Taxa-alvo do usuário primário	0,5	Coefficiente de potência do usuário secundário	0,2
SNR inicial (dB)	-20	SNR final (dB)	30
Semente para o PRNG	123	Frequência (MHz)	2000

(ξ), variando de -20 dB a 30 dB, foi escolhida de modo a abranger desde regimes severamente limitados por ruído até condições favoráveis de recepção. Além disso, a semente do gerador de números pseudoaleatórios foi fixada em 123, assegurando a reprodutibilidade dos experimentos.

A simulação também incorpora efeitos práticos típicos de sistemas sem fio, como degradações residuais de *hardware* e cancelamento imperfeito de interferência, ambos fixados em 0,05. Os coeficientes de potência atribuídos aos usuários primário e secundário foram definidos como 0,8 e 0,2, respectivamente, em conformidade com o princípio do NOMA, no qual maior potência é alocada ao usuário com condição de canal menos favorável. Por fim, a frequência de operação foi fixada em 2000 MHz, valor compatível com faixas de operação usualmente consideradas em comunicações aéreas e terrestres, afetando diretamente a perda de percurso calculada pelo modelo FSPL.

Nesta seção, analisa-se o desempenho da Probabilidade de Indisponibilidade em função da SNR média, distinguindo-se o comportamento do usuário primário e do usuário secundário no sistema NTN-NOMA. As Figuras 50, 51 e 52 apresentam a variação dessa métrica para diferentes parâmetros do modelo de canal generalizado, evidenciando a sensibilidade do sistema tanto aos parâmetros estatísticos do canal quanto às não-idealidades consideradas. Em todas as configurações avaliadas, observa-se a redução da Probabilidade de Indisponibilidade com o aumento da SNR média, comportamento coerente com a teoria de comunicações sem fio, uma vez que níveis mais elevados de SNR reduzem a probabilidade de a SNR instantânea permanecer abaixo do limiar exigido.

Ao se comparar o comportamento dos usuários, observa-se que o usuário secundário, usualmente associado à melhor condição de canal no contexto NOMA, tende a apresentar uma redução mais acentuada da Probabilidade de Indisponibilidade em determinadas faixas de SNR. Esse comportamento decorre diretamente do mecanismo de SIC: quando o usuário secundário consegue decodificar corretamente o sinal do usuário primário, ele é capaz de remover essa interferência antes da detecção do seu próprio sinal, o que resulta em melhora significativa de desempenho.

Outro aspecto relevante diz respeito à influência dos parâmetros do canal sobre as curvas de probabilidade de indisponibilidade. Em particular, o ganho de diversidade é diretamente proporcional à não linearidade do meio e à quantidade de componentes Gaussianas do modelo. Assim, à medida que α e/ou M aumentam, observa-se um aumento na inclinação assintótica

das curvas, resultando em menores valores de Probabilidade de Indisponibilidade para uma dada SNR. Esse comportamento pode ser observado nas Figuras 50a e 50b, nas quais as curvas decrescem de forma mais acentuada com o aumento da não linearidade do canal e do número de componentes Gaussianas. Ainda com base na Figura 50a, nota-se a mudança de comportamento quando o canal deixa a condição linear, correspondente a $\alpha = 2$, e passa a operar em regime não linear, isto é, para $\alpha \neq 2$.

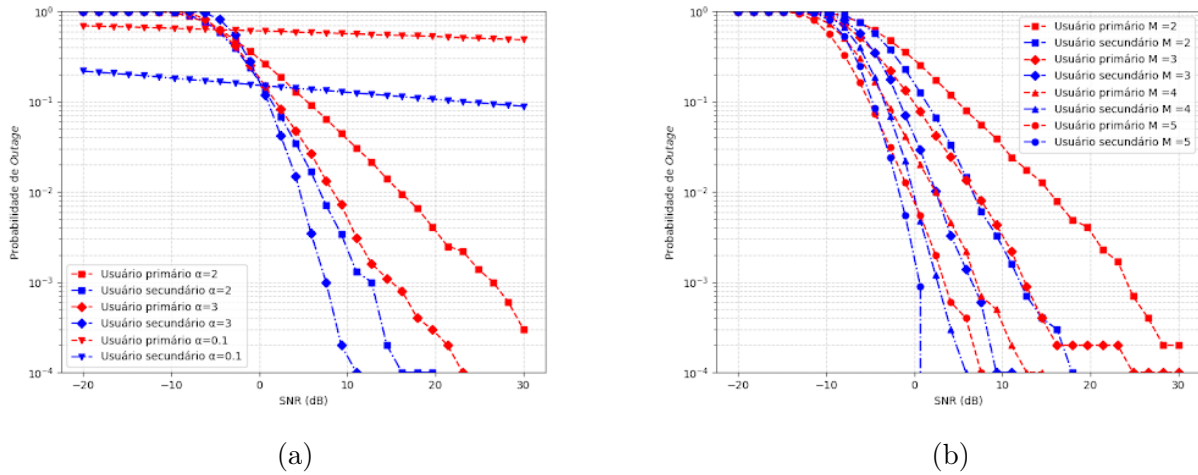


Figura 50: a) Probabilidade de Indisponibilidade (outage) *versus* a SNR para diferentes valores de α . b) Probabilidade de Indisponibilidade (outage) *versus* a SNR para diferentes valores de M .

Observa-se que o ganho de diversidade não cresce de forma ilimitada com o aumento de M . A partir dos resultados obtidos, nota-se que a variação na inclinação das curvas é mais pronunciada quando M passa de 2 para 3 do que no intervalo de 3 para 4, em concordância com o comportamento reportado em [168]. Esse resultado evidencia um efeito de retornos decrescentes, no qual os ganhos marginais de diversidade tornam-se progressivamente menores à medida que M aumenta, como ilustrado na Figura 50b.

Análises adicionais podem ser extraídas das Figuras 51a e 51b. Na Figura 51a, verifica-se que a correlação entre as componentes Gaussianas do canal impacta predominantemente o deslocamento horizontal das curvas de probabilidade de indisponibilidade, caracterizando uma variação no ganho de codificação. À medida que a correlação aumenta, reduz-se a independência entre as componentes do canal, limitando a diversidade efetivamente explorada pelo sistema e degradando o desempenho para uma mesma SNR média. Consequentemente, observa-se um deslocamento das curvas para a direita, indicando a necessidade de valores mais elevados de SNR para atingir o mesmo nível de probabilidade de indisponibilidade.

Por sua vez, na Figura 51b, o parâmetro σ , associado à potência das componentes espalhadas do canal, exerce influência significativa sobre o ganho de codificação. O aumento de σ provoca deslocamentos horizontais relevantes nas curvas, refletindo alterações nas condições de propagação. Dessa forma, embora nem a correlação nem σ alterem a inclinação assintótica das curvas, ambos impactam diretamente o posicionamento do desempenho em função da SNR, conforme também observado em [168].

Por fim, a Figura 52 evidencia o impacto do parâmetro k sobre a Probabilidade de Indisponibilidade. No contexto do modelo generalizado, esse parâmetro está associado à razão entre

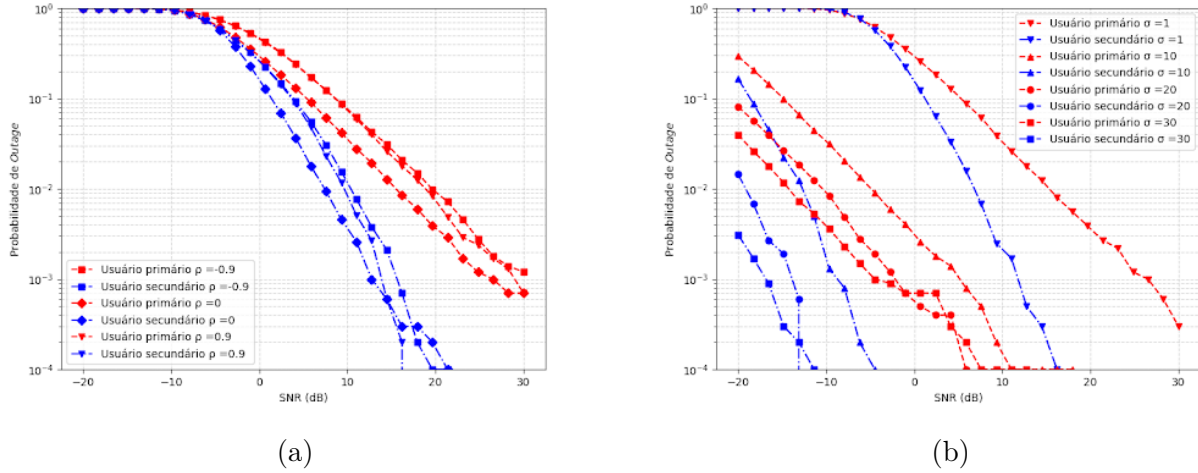


Figura 51: a) Probabilidade de Indisponibilidade (outage) *versus* a SNR para diferentes valores de correlação. b) Probabilidade de Indisponibilidade (outage) *versus* a SNR para diferentes valores de σ .

a potência das componentes especulares (LOS, por exemplo) e a potência das componentes espalhadas do canal. Observa-se que, à medida que k aumenta, as curvas de probabilidade de indisponibilidade tendem a se deslocar para a esquerda, indicando melhora no desempenho do sistema para uma mesma SNR média. Esse comportamento ocorre porque valores mais elevados de k indicam uma componente especular de alta potência (se comparada às componentes espalhadas) no canal, reduzindo a Probabilidade de Indisponibilidade.

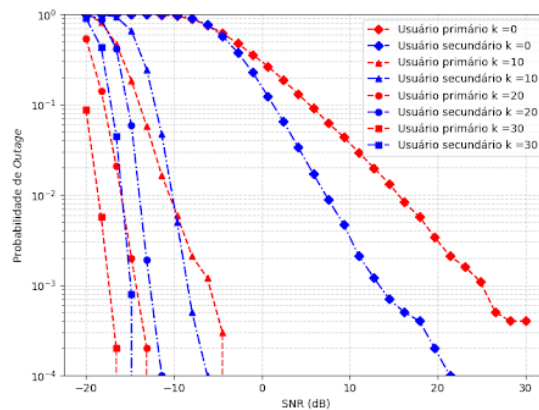


Figura 52: Probabilidade de Indisponibilidade (outage) *versus* a SNR para diferentes valores de k .

16.5 Conclusão

A crescente demanda por sistemas de comunicação mais robustos e eficientes requer que tecnologias de próxima geração operem com elevada eficiência espectral, mesmo em cenários desafiadores e heterogêneos, como aqueles associados às redes NTN. Nesse contexto, este traba-

lho avaliou o desempenho de uma rede NTN com NOMA assistida por UAV, considerando um modelo de desvanecimento generalizado capaz de representar fenômenos relevantes, incluindo não linearidades do meio, múltiplos *clusters* de propagação e correlação entre componentes do canal. As expressões analíticas obtidas para a Probabilidade de Indisponibilidade foram validadas por meio de simulações de Monte Carlo, apresentando elevada concordância e confirmando a consistência da modelagem adotada.

Os resultados obtidos indicam que o desempenho do sistema é fortemente dependente das características do canal. Em particular, verificou-se que o ganho de diversidade é diretamente influenciado pela não linearidade e pelo número de componentes Gaussianas. O aumento desses parâmetros resulta em maior inclinação da curva de Probabilidade de Indisponibilidade, o que indica maior robustez do sistema frente às variações do sinal em ambientes com múltiplos *clusters*. Adicionalmente, a análise evidenciou a relevância da estratégia de alocação de potência no contexto NOMA, demonstrando que o balanceamento adequado entre os usuários é essencial para o atendimento dos requisitos mínimos de taxa, especialmente na presença de interferência residual e limitações de *hardware*.

O modelo de canal empregado abrange tanto distribuições clássicas quanto configurações mais gerais, permitindo uma representação mais fiel das condições de propagação em cenários NTN. Como principal contribuição, este trabalho demonstra que a utilização de modelos de desvanecimento mais abrangentes aumenta a precisão da avaliação de desempenho e possibilita uma análise mais detalhada do impacto de parâmetros como não linearidade, correlação, número de componentes e potência das componentes especular e difusa na confiabilidade do sistema. Esses resultados reforçam a importância da adoção de modelos mais realistas na análise de sistemas NTN, considerando a elevada variabilidade e heterogeneidade dos enlaces envolvidos.

17 Otimização de Redes 6G NTN-NOMA em Canais de Desvanecimento Generalizado

Raphael Parreira de Souza, Bruna do Nascimento Benevides, Francisco Raimundo Albuquerque Parente e Maykon Renan Pereira da Silva
CPQD

17.1 Introdução

As *Non-Terrestrial Networks* (NTN) têm sido reconhecidos como um dos principais habilitadores dos futuros sistemas 6G, especialmente para ampliar a conectividade em áreas remotas, regiões pouco povoadas e cenários de emergência [170]. Nesse contexto, os *Unmanned Aerial Vehicles* (UAVs) têm atraído atenção significativa devido à sua mobilidade, flexibilidade de implantação e forte conectividade em *Line of Sight* (LOS). Além disso, o *Non-Orthogonal Multiple Access* (NOMA) surgiu como uma abordagem promissora para melhorar a eficiência espectral e oferecer suporte à conectividade em cenários com recursos limitados e enlaces assimétricos [171]. Ao permitir que múltiplos usuários compartilhem simultaneamente os mesmos recursos de tempo e frequência por meio da multiplexação no domínio da potência e do *Successive Interference Cancellation* (SIC), o NOMA torna-se particularmente atrativo para sistemas NTN assistidos por UAVs [172, 173].

Apesar das vantagens do NTN-NOMA, sua operação prática está sujeita a importantes limitações que podem comprometer significativamente o desempenho do sistema. Em particular, imperfeições de hardware, como não linearidades e distorções residuais, bem como falhas no SIC, podem introduzir interferência residual e reduzir os ganhos esperados em termos de taxa e confiabilidade [172, 174]. Além disso, quando o canal é representado por modelos excessivamente simplificados, a análise resultante pode não capturar adequadamente a ampla diversidade de condições de propagação encontradas em ambientes aéreos e não terrestres [175, 176]. Nesse contexto, a alocação de potência desempenha papel fundamental, pois governa diretamente o compromisso entre eficiência espectral, confiabilidade e justiça entre os usuários [177].

Neste trabalho, é estudado um sistema NTN assistido por UAV e baseado em NOMA, considerando a presença de imperfeições de hardware e erros no processo de SIC. O canal de propagação é modelado pelo desvanecimento generalizado α - η - κ - μ , permitindo representar uma ampla variedade de condições de propagação. São derivadas expressões analíticas para a probabilidade de *outage* dos usuários e investigado o impacto dos parâmetros do canal e das imperfeições práticas sobre o desempenho do sistema. Adicionalmente, uma estratégia de alocação de potência baseada no algoritmo *Particle Swarm Optimization* (PSO) é empregada para minimizar a probabilidade de *outage*, possibilitando a melhoria da confiabilidade da comunicação.

17.2 Modelo do Sistema

Nesta seção, descrevemos a arquitetura da rede e os modelos matemáticos de *Signal to Interference plus Noise Ratio* (SINR) para um sistema downlink UAV-NOMA, conforme ilustrado na Fig.53. Com base no cenário apresentado em [169], o modelo caracteriza um UAV atendendo múltiplos usuários terrestres, incorporando os efeitos da alocação de potência, imperfeições de hardware e limitações do SIC.

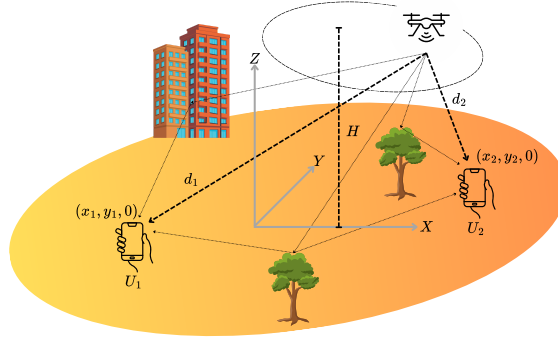


Figura 53: Modelo do sistema NTN.

17.2.1 Modelo de Desvanecimento

Em ambientes multipercurso, o sinal recebido no receptor é composto pela soma de diversas componentes propagadas, cada uma sujeita a amplitudes e fases aleatórias. Devido ao *Central Limit Theorem* (CLT), a soma de múltiplas componentes independentes tende a apresentar distribuição Gaussiana. Com base nessa propriedade, considera-se neste trabalho um modelo de desvanecimento generalizado capaz de representar múltiplos agrupamentos de espalhamento, componentes especulares e efeitos não lineares presentes em sistemas NTN- NOMA.

A potência do canal, denotada por B , é definida como a soma de M variáveis aleatórias Gaussianas quadráticas, dada por:

$$B^{\frac{\alpha}{2}} = \sum_{i=1}^M X_i^2, \quad (211)$$

em que $\alpha > 0$ representa o parâmetro de não linearidade do canal, enquanto M corresponde ao número de componentes Gaussianas presentes no ambiente de propagação.

O vetor aleatório

$$X \triangleq [X_1, X_2, \dots, X_M]^T \quad (212)$$

segue uma Função Densidade de Probabilidade (*Probability Density Function* (PDF)) Gaussiana multivariada com vetor média

$$m \triangleq E[X] = [E[X_1], E[X_2], \dots, E[X_M]]^T \quad (213)$$

e matriz de covariância definida como

$$\Sigma \triangleq E[(X - m)(X - m)^T], \quad (214)$$

em que

$$E[X_i] = m_i \quad (215)$$

e

$$V[X_i] = \sigma_i^2, \quad \forall i \in M. \quad (216)$$

A matriz de covariância positiva definida pode ser escrita como

$$\Sigma = \begin{bmatrix} \sigma_1^2 & \rho_{1,2}\sigma_1\sigma_2 & \cdots & \rho_{1,M}\sigma_1\sigma_M \\ \rho_{1,2}\sigma_1\sigma_2 & \sigma_2^2 & \cdots & \rho_{2,M}\sigma_2\sigma_M \\ \vdots & \vdots & \ddots & \vdots \\ \rho_{1,M}\sigma_1\sigma_M & \rho_{2,M}\sigma_2\sigma_M & \cdots & \sigma_M^2 \end{bmatrix} \quad (217)$$

em que

$$\rho_{i,j} \equiv \rho(X_i, X_j) \quad (218)$$

representa o coeficiente de correlação entre as componentes Gaussianas X_i e X_j .

Consequentemente, a PDF Gaussiana conjunta de X é dada por:

$$f_X(x) = \frac{\exp \left[-\frac{1}{2}(x - m)^T \Sigma^{-1}(x - m) \right]}{[(2\pi)^M \det(\Sigma)]^{1/2}} \quad (219)$$

com

$$x \triangleq [x_1, x_2, \dots, x_M]^T. \quad (220)$$

A análise exata do desempenho do sistema em termos de probabilidade de outage para esse modelo de desvanecimento depende do cálculo de uma integral multidimensional sobre a PDF Gaussiana multivariada apresentada em (219).

A PDF da potência do canal associada ao modelo de desvanecimento da Equação (211) pode ser expressa como:

$$f_B(\beta) = \alpha \beta^{\frac{\alpha}{2}-1} [(2^{M+3}\pi)^M \det(\Sigma)]^{1/2} \int_0^{\beta^{\alpha/2}} \cdots \int_0^{\beta^{\alpha/2}} \left(\prod_{i=1}^M \sqrt{s_i} \right)^{-1} \quad (221)$$

$$\times \sum_{j=1}^{2^M} \exp \left[-\frac{1}{2}(r_j - m)^T \Sigma^{-1}(r_j - m) \right] ds_2 \cdots ds_M, \quad (222)$$

em que $s_i \triangleq x_i^2$ e

$$s \triangleq [s_1, s_2, \dots, s_M]^T, \quad (223)$$

com

$$s_1 = \beta^{\alpha/2} - \sum_{i=2}^M s_i. \quad (224)$$

Além disso, o vetor r_j é definido para cada $j \in \{1, \dots, 2^M\}$ considerando todas as possíveis combinações das componentes $\{\sqrt{s_i}\}_{i=1}^M$.

A partir dessa formulação, a probabilidade de outage (P_{out}), definida como a probabilidade de que a *Signal-to-Noise Ratio* (SNR) instantânea γ fique abaixo de um limiar requerido γ_{th} , é dada por:

$$P_{out} \triangleq \int_0^{\frac{\gamma_{th}}{\bar{\gamma}}} f_B(\beta) d\beta, \quad (225)$$

em que $\bar{\gamma}$ representa a SNR média considerando $E[B] = 1$.

Essa métrica também pode ser analisada em termos de ganhos de diversidade e de codificação. Conforme discutido ao longo deste trabalho, tais ganhos são diretamente influenciados pelos diferentes parâmetros do modelo de desvanecimento generalizado considerado.

17.2.2 Geometria do UAV e Caracterização do Canal

Considera-se um UAV posicionado a uma altitude fixa H , com coordenadas horizontais (x, y) . Os usuários terrestres encontram-se distribuídos dentro da área de cobertura, sendo o m -ésimo usuário localizado em $(x_m, y_m, 0)$.

A distância Euclidiana entre o UAV e o usuário m é definida como:

$$d_m = \sqrt{(x_m - x)^2 + (y_m - y)^2 + H^2}. \quad (226)$$

Para modelar a atenuação do sinal em larga escala, adota-se o modelo de *Free Space Path Loss* (FSPL), altamente adequado para enlaces *Air-to-Ground* (A2G) com predominância de componente LOS. A perda de percurso em dB é dada por:

$$PL_{FSPL}(dB) = 32.4 + 20 \log_{10}(f_{MHz}) + 20 \log_{10}(d_{km}), \quad (227)$$

em que f_{MHz} representa a frequência da portadora em MHz e d_{km} corresponde à distância em quilômetros.

Para fins analíticos, a perda de percurso é convertida para escala linear através de

$$L(d) = 10^{-\frac{PL_{FSPL}(dB)}{10}}. \quad (228)$$

Define-se ainda

$$L_0 = 10^{-\frac{PL_0(dB)}{10}} \quad (229)$$

como sendo o termo de normalização para uma distância de referência de 1 km, em que:

$$PL_0(dB) = 32.4 + 20 \log_{10}(f_{MHz}). \quad (230)$$

Assim, o ganho efetivo de potência do canal para um determinado usuário é modelado como

$$|h|^2 = \frac{L(d)}{L_0} B \quad (231)$$

em que B representa a potência do canal associada ao desvanecimento de pequena escala.

Com base nesses ganhos, distingue-se um canal primário, associado ao menor ganho, e um canal secundário, associado ao maior ganho.

17.3 Transmissão NOMA e Codificação por Superposição

Diferentemente dos métodos convencionais de acesso múltiplo ortogonal, o NOMA permite que múltiplos usuários compartilhem simultaneamente os mesmos recursos de frequência e tempo por meio da multiplexação no domínio da potência.

No transmissor do UAV, as mensagens destinadas aos usuários s_1 e s_2 são sobrepostas em um único sinal x utilizando Codificação por Superposição (*Superposition Coding* (SC)):

$$x = \sqrt{p_1 P} s_1 + \sqrt{p_2 P} s_2, \quad (232)$$

em que P representa a potência total de transmissão, enquanto p_1 e p_2 são os coeficientes de alocação de potência.

Recepção do Sinal e Análise da SINR

O sinal recebido pelo usuário k está sujeito às condições individuais do canal h_k e ao *Additive White Gaussian Noise* (AWGN) n_k , sendo expresso como

$$y_k = h_k^H x + n_k. \quad (233)$$

Substituindo (232) em (233), obtém-se:

$$y_k = h_k^H \left(\sqrt{p_1 P} s_1 + \sqrt{p_2 P} s_2 \right) + n_k. \quad (234)$$

Em sistemas práticos, o desempenho é degradado por imperfeições de hardware e pela incapacidade de realizar cancelamento perfeito de interferência.

Define-se:

$$\xi = \frac{P}{N_0} \quad (235)$$

como sendo a razão sinal-ruído de transmissão (SNR), onde N_0 corresponde à densidade espectral de potência do ruído.

Para o usuário distante, o sinal do Usuário 2 é tratado como interferência. Considerando imperfeições de hardware representadas por κ , com $0 \leq \kappa \leq 1$, a SINR do Usuário 1 é dada por:

$$\gamma_1 = \frac{\xi |h_1|^2 p_1}{\xi |h_1|^2 (p_2 + \kappa^2) + 1}. \quad (236)$$

O usuário próximo utiliza SIC para remover o sinal do Usuário 1 antes de decodificar sua própria informação. Entretanto, devido a limitações práticas, uma parcela residual de interferência persiste após o cancelamento.

Assim, a SINR do Usuário 2 é expressa como

$$\gamma_2 = \frac{\xi |h_2|^2 p_2}{\xi |h_2|^2 (p_1 \epsilon + \kappa^2) + 1} \quad (237)$$

em que $\epsilon \in [0,1]$ representa a imperfeição do SIC. Quando $\epsilon = 0$, tem-se SIC perfeito, enquanto $\epsilon = 1$ indica que toda a interferência do Usuário 1 permanece sem cancelamento.

17.4 Estratégia de Otimização da Alocação de Potência

Esta seção apresenta uma estratégia metaheurística para a alocação ótima de potência em sistemas UAV-NOMA. Como as funções objetivo derivadas da SINR possuem natureza não convexa, tornando os métodos convencionais baseados em gradiente ineficazes, uma abordagem metaheurística é adotada para determinar os coeficientes ótimos de potência e melhorar o desempenho do sistema.

17.4.1 Formulação do Problema

A estratégia Max-Min Rate garante equidade ao maximizar a menor taxa entre U_1 e U_2 . Essa abordagem prioriza o usuário limitante do sistema, promovendo equilíbrio entre as taxas e reduzindo a probabilidade global de outage.

O problema de otimização é formalmente definido como:

$$\max_{p_1, p_2} \min(R_1, R_2) \quad (238)$$

sujeito às seguintes restrições:

$$0 < p_1 < 1 \quad (239)$$

$$p_1 + p_2 = 1, \quad (240)$$

em que p_1 e p_2 representam os coeficientes de alocação de potência, enquanto R_1 e R_2 incorporam imperfeições de hardware e SIC residual, conforme definido nas seguintes expressões de taxa:

$$R_1 = \log_2 \left(1 + \frac{\xi |h_1|^2 p_1}{\xi |h_1|^2 (p_2 + \kappa^2) + 1} \right) \quad (241)$$

$$R_2 = \log_2 \left(1 + \frac{\xi |h_2|^2 p_2}{\xi |h_2|^2 (p_1 + \kappa^2) + 1} \right). \quad (242)$$

17.4.2 Algoritmo de Otimização por Enxame de Partículas

Para resolver o problema não convexo apresentado em (238), emprega-se o algoritmo Particle Swarm Optimization (PSO). Essa metaheurística baseada em população, inspirada no comportamento social coletivo, é adequada para explorar superfícies não convexas sem a necessidade de informações de gradiente. A implementação é descrita a seguir:

1. **Inicialização:** Uma população contendo N_p partículas é criada. A posição de cada partícula corresponde a um valor candidato para p_1 , inicializado aleatoriamente. Cada partícula também recebe uma velocidade inicial aleatória (v_i).
2. **Avaliação Inicial:** A função de aptidão ($f_{p,i}$), definida como $\min(R_1, R_2)$, é calculada para cada partícula. Em seguida, são armazenadas a melhor posição individual de cada partícula ($p_{best,i}$) e a melhor posição global do enxame (g_{best}).
3. **Atualização Iterativa:** Durante um número máximo de iterações I_{max} , a velocidade e a posição de cada partícula são atualizadas. A atualização da velocidade é influenciada por:
 - sua velocidade atual (componente de inércia, ponderada por w);
 - sua melhor posição individual (componente cognitiva, ponderada por c_1);
 - a melhor posição global do enxame (componente social, ponderada por c_2).
4. **Limites e Avaliação:** A nova posição é limitada ao intervalo $[0.001, 0.999]$. Em seguida, o valor de $f_{p,i}$ é recalculado. Caso o novo valor de aptidão seja superior ao anterior, atualizam-se $p_{best,i}$ e, potencialmente, g_{best} .
5. **Parada Antecipada:** Para aumentar a eficiência computacional, foi incorporado um mecanismo de parada antecipada. Esse mecanismo utiliza um parâmetro de tolerância (τ) e um parâmetro de paciência (P). Caso a melhoria da melhor solução global permaneça inferior a τ durante P iterações consecutivas, o processo de otimização é encerrado.
6. **Resultado:** Após a última iteração, ou quando a condição de parada antecipada é satisfeita, a posição g_{best} é retornada como solução ótima do PSO.

Algorithm 17.1 Algoritmo PSO com Parada Antecipada

```

1: Inicialização: Definir  $|h_1|^2, |h_2|^2, \xi, \kappa, \epsilon$  e os parâmetros do PSO:  $N_p, I_{max}, w, c_1, c_2, \tau$  e  $P$ .
2: Inicializar  $N_p$  partículas com posições aleatórias  $x_i \in [0.001, 0.999]$  e velocidades  $v_i$ .
3: Definir  $p_{best,i} \leftarrow x_i$  e avaliar  $f_{p_{best,i}}$ .
4: Definir  $g_{best}$  e avaliar  $f_{g_{best}}$ .
5:  $stagnation\_counter \leftarrow 0$ .
6: for  $iter = 1$  to  $I_{max}$  do
7:    $previous\_f_{g_{best}} \leftarrow f_{g_{best}}$ .
8:   for  $i = 1$  to  $N_p$  do
9:     Gerar números aleatórios  $r_1, r_2 \in [0, 1]$ .
10:    Atualizar:  $v_i \leftarrow w \cdot v_i + c_1 r_1 (p_{best,i} - x_i) + c_2 r_2 (g_{best} - x_i)$ .
11:    Atualizar:  $x_i \leftarrow x_i + v_i$ .
12:    Limitar  $x_i$  ao intervalo  $[0.001, 0.999]$ .
13:     $f_{current} \leftarrow f_{p,i}$ .
14:    if  $f_{current} < f_{p_{best,i}}$  then
15:       $p_{best,i} \leftarrow x_i$  e  $f_{p_{best,i}} \leftarrow f_{current}$ .
16:      if  $f_{current} < f_{g_{best}}$  then
17:         $g_{best} \leftarrow x_i$  e  $f_{g_{best}} \leftarrow f_{current}$ .
18:      end if
19:    end if
20:  end for
21:  if  $|f_{g_{best}} - previous\_f_{g_{best}}| / previous\_f_{g_{best}} < \tau$  then
22:     $stagnation\_counter \leftarrow stagnation\_counter + 1$ .
23:  else
24:     $stagnation\_counter \leftarrow 0$ .
25:  end if
26:  if  $stagnation\_counter \geq P$  then
27:    break.
28:  end if
29: end for
30: Saída: Coeficiente ótimo de alocação de potência  $g_{best}$ .
  
```

Para validar os ganhos teóricos proporcionados pela estratégia Max-Min baseada em PSO, a próxima seção apresenta uma avaliação numérica abrangente. O desempenho do sistema é analisado sob diferentes regimes de SNR, comparando a solução otimizada com abordagens convencionais, a fim de quantificar as melhorias em termos de equidade e confiabilidade.

17.5 Resultados e Discussões

Apresenta-se agora uma análise comparativa detalhada dos resultados de simulação obtidos sob esquemas de alocação de potência arbitrária e baseada em PSO. O desempenho dessas estratégias é avaliado em termos da taxa média alcançável e da probabilidade de outage em função da SNR. Essa análise quantifica os ganhos proporcionados pelo processo de otimização e demonstra a eficácia do PSO como uma solução de baixa complexidade, escalável e eficiente para alocação de potência.

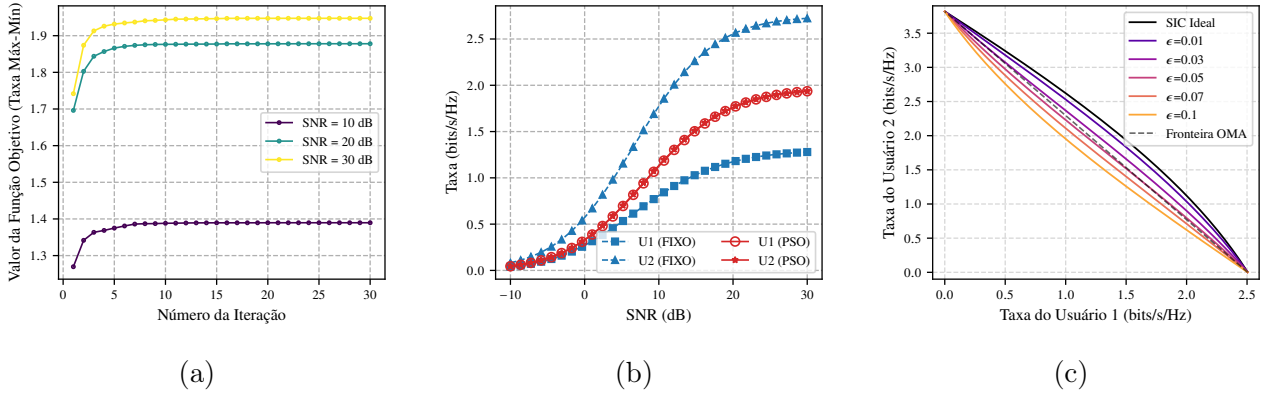


Figura 54: (a) Curva de convergência por iteração. (b) Comparação da taxa de transmissão com alocação de potência fixa e PSO. (c) Região de taxas considerando os limites de NOMA e OMA.

A Fig. 54(a) demonstra a eficiência do algoritmo PSO, alcançando consistentemente a estabilidade entre 10 e 15 iterações. Essa rápida convergência mantém-se robusta para níveis de SNR de 10, 20 e 30 dB, confirmando a adequação do algoritmo para sistemas de rádio em tempo real, nos quais valores mais elevados de SNR resultam previsivelmente em melhores valores da função objetivo.

A Fig. 54(b) mostra que o algoritmo PSO obtém um equilíbrio quase ótimo ao priorizar a equidade entre as taxas do sistema. Enquanto a alocação fixa resulta em uma diferença significativa entre as taxas dos usuários, o PSO conduz as taxas do Usuário 1 e do Usuário 2 à convergência e sobreposição. Ao redistribuir recursos do usuário mais forte para aquele com menor SINR, o algoritmo assegura um desempenho mais equilibrado entre ambos.

Uma observação importante nos resultados é a saturação da taxa alcançável em altos níveis de SNR, particularmente quando a SNR se aproxima de 30 dB. Ao analisar as expressões de SINR para os Usuários 1 e 2 à medida que o termo de ruído $\frac{1}{\xi \cdot h^2}$ tende a zero, obtêm-se os seguintes limites:

$$\text{SINR}_1 = \frac{p}{(1-p) + \kappa^2 + \frac{1}{\xi \cdot h_1^2}} \xrightarrow{\xi \rightarrow \infty} \frac{p}{(1-p) + \kappa^2} \quad (243)$$

$$\text{SINR}_2 = \frac{1-p}{\epsilon \cdot p + \kappa^2 + \frac{1}{\xi \cdot h_2^2}} \xrightarrow{\xi \rightarrow \infty} \frac{1-p}{\epsilon \cdot p + \kappa^2}. \quad (244)$$

Essas derivações confirmam que o desempenho do sistema é estritamente limitado pelo fator de imperfeição de hardware (κ) e pelo erro residual de SIC (ϵ). Esses parâmetros estabelecem um limite físico para o desempenho, tornando ineficazes aumentos adicionais na SNR ou no ganho de canal, uma vez que o sistema passa a ser limitado predominantemente por interferências e imperfeições.

A Fig. 54(c) ilustra a região de taxas (Usuário 2 versus Usuário 1), definindo o limite em que a vantagem de eficiência espectral do NOMA em relação ao OMA desaparece. Enquanto o SIC ideal ($\epsilon = 0$) produz uma região convexa que supera consistentemente o OMA, o aumento de ϵ (de 0.01 para 0.1) reduz essa região em direção à origem.

Observa-se ainda que, para erros elevados de SIC (por exemplo, $\epsilon \geq 0.07$), o NOMA torna-se “degradado” e passa a apresentar desempenho inferior ao limite do OMA. Nesse regime, a penalidade causada pela interferência residual supera o ganho proporcionado pela multiplexação

no domínio da potência, tornando a separação ortogonal mais eficiente do que o cancelamento imperfeito de interferência. Assim, a viabilidade do NOMA depende diretamente da precisão do SIC. Acima de determinado limiar de ϵ , o OMA tradicional passa a ser a alternativa mais eficiente para evitar o acúmulo de efeitos interferentes.

A probabilidade de outage em função da SNR média para diferentes parâmetros do modelo de canal generalizado é apresentada na Fig. 55. Essa figura também inclui o desempenho do sistema otimizado por meio do algoritmo PSO, permitindo uma comparação direta com as configurações não otimizadas. Os resultados destacam a sensibilidade do sistema tanto às características estatísticas do canal quanto às condições não ideais consideradas. Como níveis mais elevados de SNR média reduzem a probabilidade de a SNR instantânea permanecer abaixo do limiar requerido, observa-se, em todas as configurações avaliadas, uma redução da probabilidade de outage com o aumento da SNR média, conforme esperado.

Na Fig. 55(a), é demonstrado o impacto do parâmetro de correlação ρ . Observa-se que a correlação influencia principalmente o deslocamento horizontal das curvas de outage, o que está associado a variações no ganho de codificação. À medida que ρ aumenta, os componentes Gaussianos tornam-se mais dependentes, reduzindo a diversidade efetiva disponível no sistema. Consequentemente, as curvas deslocam-se para regiões de maior SNR, indicando que níveis mais elevados de SNR são necessários para atingir a mesma probabilidade de outage.

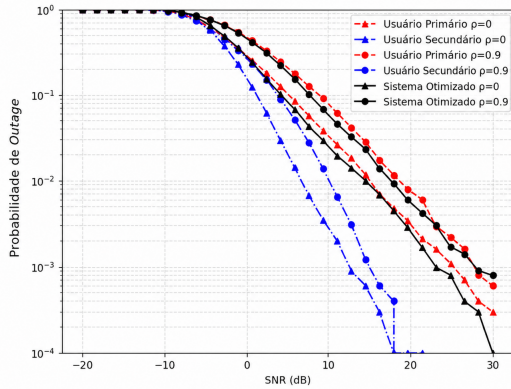
Na Fig. 55(b), observa-se que o aumento de M resulta em melhora do desempenho de outage devido à diversidade adicional introduzida por um maior número de componentes Gaussianos. Entretanto, essa melhora não ocorre de forma linear, pois o ganho mais significativo acontece na transição de $M = 2$ para $M = 3$, enquanto incrementos adicionais proporcionam benefícios progressivamente menores, indicando retornos decrescentes de diversidade.

Conforme apresentado na Fig. 55(c), o parâmetro σ , associado à intensidade das componentes espalhadas do modelo, exerce influência significativa sobre o ganho de codificação. O aumento do desvio padrão provoca um deslocamento horizontal das curvas, refletindo alterações nas condições do canal. Embora σ e o parâmetro de correlação não modifiquem a inclinação assintótica das curvas, ambos desempenham papel fundamental na determinação dos níveis de SNR necessários para uma determinada probabilidade de outage.

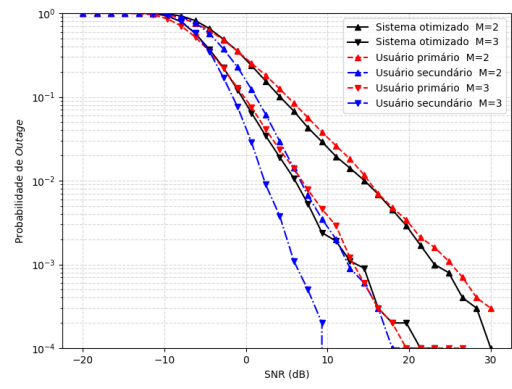
Na Fig. 55(d), o parâmetro k representa a razão entre a potência das componentes dominantes e a potência das componentes espalhadas. O aumento de k desloca as curvas para menores valores de SNR, indicando melhora no desempenho do sistema, já que a mesma probabilidade de outage pode ser alcançada com menor potência de transmissão.

Por fim, a Fig. 55(e) apresenta o impacto do parâmetro de não linearidade α . Valores mais elevados de α produzem um decaimento mais acentuado das curvas de outage, refletindo aumento no ganho de diversidade e na confiabilidade global do sistema. Além disso, observa-se uma mudança significativa de comportamento na transição do regime linear ($\alpha = 2$) para condições não lineares ($\alpha \neq 2$).

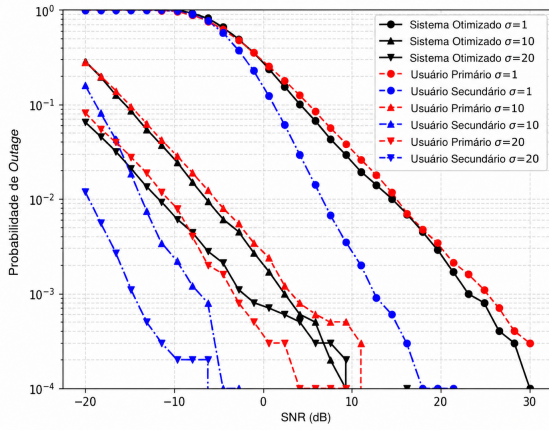
De forma geral, os resultados indicam que os parâmetros do canal generalizado governam simultaneamente os ganhos de diversidade e de codificação. A otimização max-min baseada em PSO melhora a confiabilidade global ao equilibrar os enlaces de comunicação. Embora o usuário mais forte sofra uma penalização de desempenho, o ganho obtido pelo usuário mais fraco reduz a probabilidade total de outage. Essa estratégia assegura um desempenho robusto mesmo na presença de imperfeições do canal e efeitos de não linearidade.



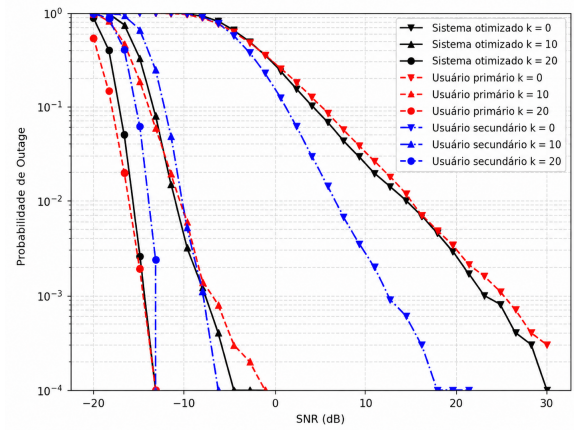
(a)



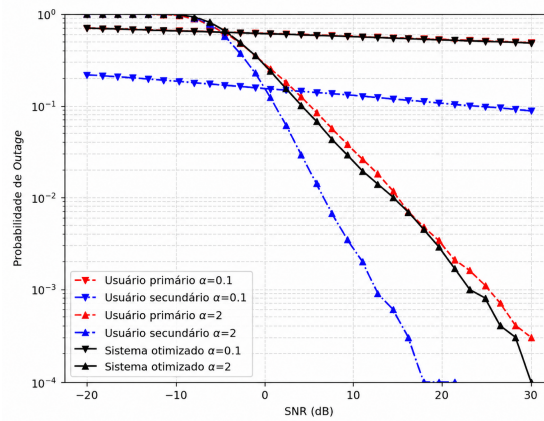
(b)



(c)



(d)



(e)

Figura 55: Probabilidade de outage em função da SNR média para diferentes configurações dos parâmetros ρ , M , σ , k e α , incluindo a curva correspondente ao método de otimização proposto.

17.6 Conclusão

Neste trabalho, foi avaliado o desempenho de sistemas NTN-NOMA assistidos por UAVs sob um modelo de desvanecimento generalizado, considerando também imperfeições práticas de hardware e limitações do SIC. Diferentemente de abordagens simplificadas encontradas na literatura, o modelo proposto incorporou simultaneamente efeitos de não linearidade do canal, componentes correlacionadas de desvanecimento, perda de percurso em espaço livre e interferência residual, permitindo uma análise mais realista de cenários NTN voltados para aplicações 6G.

A partir da formulação de um problema de otimização baseado no critério Max-Min, foi aplicada uma estratégia de alocação de potência utilizando o algoritmo PSO. Os resultados obtidos demonstraram que a abordagem proposta é capaz de melhorar significativamente a justiça entre os usuários, redistribuindo potência de forma adaptativa para favorecer os enlaces mais degradados. Como consequência, observou-se uma redução relevante da probabilidade de outage para usuários com menor qualidade de canal, além de uma distribuição de taxas mais equilibrada quando comparada a esquemas convencionais de alocação fixa.

Os resultados também evidenciaram que o algoritmo PSO apresenta elevada eficiência computacional, convergindo rapidamente em poucas iterações, o que o torna adequado para aplicações em tempo real em sistemas 6G. Além disso, verificou-se que parâmetros associados ao desvanecimento generalizado, como o número de componentes Gaussianas M , o parâmetro de correlação ρ , o fator k e o parâmetro de não linearidade α , exercem impacto direto sobre os ganhos de diversidade e codificação, influenciando fortemente o desempenho do sistema em termos de confiabilidade.

Outro aspecto importante observado foi a existência de um limite de desempenho em regimes de alta SNR. Nessa condição, os efeitos das imperfeições de hardware e do SIC residual tornam-se dominantes, impedindo ganhos adicionais mesmo com o aumento da potência de transmissão. Os resultados mostraram ainda que, acima de determinados níveis de interferência residual, o desempenho do NOMA pode se degradar a ponto de soluções ortogonais, como *Orthogonal Multiple Access* (OMA), apresentarem desempenho superior. Dessa forma, a precisão do SIC torna-se um requisito fundamental para a viabilidade prática de sistemas NOMA em cenários NTN.

De maneira geral, os resultados obtidos fornecem importantes diretrizes para o projeto de futuras redes 6G baseadas em NTN, especialmente em relação ao impacto das não linearidades do canal, das limitações práticas de hardware e das estratégias inteligentes de alocação de potência sobre a confiabilidade e eficiência espectral do sistema.

Como perspectivas futuras, pretende-se investigar estratégias conjuntas de otimização envolvendo posicionamento dinâmico de UAVs, formação de feixes, técnicas avançadas de precodificação e algoritmos híbridos de aprendizado de máquina para adaptação em tempo real das políticas de alocação de potência em ambientes altamente dinâmicos.

18 Conclusão Geral

Este relatório técnico apresenta os resultados obtidos no âmbito da Atividade 2.2 do Projeto Brasil 6G - Fase 3, referentes à concepção da rede de acesso para sistemas 6G. Foram desenvolvidos estudos em diversas frentes de pesquisa, abrangendo aspectos teóricos, metodológicos e aplicados da camada física e de acesso ao meio. As contribuições trataram de temas como transferência de energia sem fio assistida por RIS, pré-distorção digital em sistemas ópticos-sem fio integrados, redução de PAPR via aprendizado por reforço, comunicação orientada a tarefas e semântica, protocolos de acesso baseados em *Age of Information*, segurança contra estações base falsas, técnicas de múltiplo acesso para IoT massivo e comunicação via satélite, otimização de recursos em sistemas *cell-free massive* MIMO, seleção de feixes assistida por aprendizado de máquina, e alocação de potência em redes RSMA e NOMA terrestres, aéreas e não terrestres.

De modo geral, as técnicas investigadas buscaram atender aos requisitos previstos para os sistemas 6G, como maior eficiência espectral e energética, maior confiabilidade e justiça entre os usuários, além de soluções voltadas à escalabilidade, segurança e conectividade em cenários heterogêneos. Destaca-se, ainda, o uso crescente de técnicas de otimização e de algoritmos de aprendizado de máquina como ferramentas para lidar com a complexidade e a dinâmica desses novos ambientes de comunicação.

Referências

- [1] B. S. de C. da Silva, L. A. M. Pereira, and L. L. Mendes, “Reinforcement learning to reduce papr in communication systems,” in *2025 Workshop on Communication Networks and Power Systems (WCNPS)*, 2025, pp. 1–6.
- [2] F. A. Tondo, J. M. de Souza Sant’Ana, S. Montejo-Sánchez, O. L. A. López, S. Céspedes, and R. D. Souza, “Non-Orthogonal Multiple-Access Strategies for Direct-to-Satellite IoT Networks,” 2024. [Online]. Available: <https://arxiv.org/abs/2409.02748>
- [3] B. Salehi, G. Reus-Muns, D. Roy, Z. Wang, T. Jian, J. Dy, S. Ioannidis, and K. Chowdhury, “Deep learning on multimodal sensor data at the wireless edge for vehicular network,” *IEEE Transactions on Vehicular Technology*, vol. 71, no. 7, pp. 7639–7655, 2022.
- [4] J. Ruseckas, G. Molis, and H. Bogucka, “Mimo beam selection in 5g using neural networks,” *International Journal of Electronics and Telecommunications*, vol. 67, no. 4, 2021.
- [5] J. C. Manjarres, D. O. Cardoso, and J. F. De Rezende, “A low-complexity window-based learning framework for adaptive 6g beam selection using wisard,” *Computer Networks*, p. 112128, 2026.
- [6] O. M. Rosabal, O. L. A. López, H. Alves, and M. Latva-Aho, “Sustainable RF Wireless Energy Transfer for Massive IoT: Enablers and Challenges,” *IEEE Access*, vol. 11, pp. 133 979–133 992, 2023.
- [7] V. D. P. Souto, O. M. Rosabal, S. Montejo-Sánchez, O. L. A. López, H. Alves, and R. D. Souza, “An Analog Beamforming Strategy Based on Statistical Channel Information for Recharging Wireless Powered Sensor Networks,” *IEEE Sensors Journal*, vol. 24, no. 14, pp. 23 025–23 033, 2024.
- [8] J. Shi, P. Cong, L. Zhao, X. Wang, S. Wan, and M. Guizani, “A Two-Stage Strategy for UAV-Enabled Wireless Power Transfer in Unknown Environments,” *IEEE Trans. on Mobile Computing*, vol. 23, no. 2, pp. 1785–1802, 2024.
- [9] Q. Wu, S. Zhang, B. Zheng, C. You, and R. Zhang, “Intelligent Reflecting Surface-Aided Wireless Communications: A Tutorial,” *IEEE Trans. on Communications*, vol. 69, no. 5, pp. 3313–3351, 2021.
- [10] H. Yang, X. Yuan, J. Fang, and Y.-C. Liang, “Reconfigurable intelligent surface aided constant-envelope wireless power transfer,” *IEEE Trans. on Signal Processing*, vol. 69, pp. 1347–1361, 2021.
- [11] X. Yu, Q. Sun, X. Chen, J. Zhang, C. Xu, and Y. Yang, “Spectral Efficiency Analysis for RISs-Aided Wireless-Powered Cell-Free Massive MIMO,” in *IEEE Wireless Communications and Networking Conference (WCNC)*, 2024, pp. 1–6.
- [12] L. Zhao, Z. Wang, and X. Wang, “Wireless Power Transfer Empowered by Reconfigurable Intelligent Surfaces,” *IEEE Systems Journal*, vol. 15, no. 2, pp. 2121–2124, 2021.

- [13] S. Shen, B. Clerckx, and R. Murch, “Modeling and Architecture Design of Reconfigurable Intelligent Surfaces Using Scattering Parameter Network Analysis,” *IEEE Trans. on Wireless Communications*, vol. 21, no. 2, pp. 1229–1243, 2022.
- [14] H. Li, S. Shen, M. Nerini, and B. Clerckx, “Reconfigurable Intelligent Surfaces 2.0: Beyond Diagonal Phase Shift Matrices,” *IEEE Communications Magazine*, vol. 62, no. 3, pp. 102–108, 2024.
- [15] M. Nerini, S. Shen, and B. Clerckx, “Closed-Form Global Optimization of Beyond Diagonal Reconfigurable Intelligent Surfaces,” *IEEE Trans. on Wireless Communications*, vol. 23, no. 2, pp. 1037–1051, 2024.
- [16] H. Li, S. Shen, and B. Clerckx, “A Dynamic Grouping Strategy for Beyond Diagonal Reconfigurable Intelligent Surfaces With Hybrid Transmitting and Reflecting Mode,” *IEEE Trans. on Vehicular Technology*, vol. 72, no. 12, pp. 16 748–16 753, 2023.
- [17] M. Nerini, S. Shen, and B. Clerckx, “Discrete-Value Group and Fully Connected Architectures for Beyond Diagonal Reconfigurable Intelligent Surfaces,” *IEEE Trans. on Vehicular Technology*, vol. 72, no. 12, pp. 16 354–16 368, 2023.
- [18] Y. Zhou, Y. Liu, H. Li, Q. Wu, S. Shen, and B. Clerckx, “Optimizing Power Consumption, Energy Efficiency, and Sum-Rate Using Beyond Diagonal RIS—A Unified Approach,” *IEEE Trans. on Wireless Communications*, vol. 23, no. 7, pp. 7423–7438, 2024.
- [19] H. Li, S. Shen, and B. Clerckx, “Synergizing Beyond Diagonal Reconfigurable Intelligent Surface and Rate-Splitting Multiple Access,” *IEEE Trans. on Wireless Communications*, vol. 23, no. 8, pp. 8717–8729, 2024.
- [20] I. Ahmed, H. Khammari, A. Shahid, A. Musa, K. S. Kim, E. De Poorter, and I. Moerman, “A Survey on Hybrid Beamforming Techniques in 5G: Architecture and System Model Perspectives,” *IEEE Communications Surveys & Tutorials*, vol. 20, no. 4, pp. 3060–3097, 2018.
- [21] V. D. P. Souto, O. M. Rosabal, S. Montejó-Sánchez, O. L. A. López, H. Alves, and R. D. Souza, “An analog beamforming strategy based on statistical channel information for recharging wireless powered sensor networks,” *IEEE Sensors Journal*, pp. 1–1, 2024.
- [22] E. Boshkovska, D. W. K. Ng, N. Zlatanov, and R. Schober, “Practical Non-Linear Energy Harvesting Model and Resource Allocation for SWIPT Systems,” *IEEE Communications Letters*, vol. 19, no. 12, pp. 2082–2085, 2015.
- [23] Powercast, “Lifetime Power® Energy Harvesting Development Kit for Battery Recharging,” <https://www.powercastco.com/p1110-eval-01-gate/>, accessed: 2024-16-02.
- [24] J. Kennedy and R. Eberhart, “Particle swarm optimization,” in *Proceedings of International Conference on Neural Networks*, vol. 4, Nov 1995, pp. 1942–1948 vol.4.
- [25] Q. Pham, D. C. Nguyen, S. Mirjalili, D. T. Hoang, D. N. Nguyen, P. N. Pathirana, and W. Hwang, “Swarm Intelligence for Next-Generation Wireless Networks: Recent Advances and Applications,” *CoRR*, vol. abs/2007.15221, 2020. [Online]. Available: <https://arxiv.org/abs/2007.15221>

- [26] M. Mello, L. Mendes, and T. Barbosa, "Spectrum Efficient GFDM Based on Faster Than Nyquist Signaling," *JCIS*, vol. 35, no. 1, pp. 349–356, Dec 2020.
- [27] L. L. Mendes, N. Michailow, M. Matth  , I. Gaspar, D. Zhang, and G. Fettweis, "GFDM: providing flexibility for the 5G physical layer," *Opportunities in 5G Networks: A Research and Development Perspective*, p. 325, 2016.
- [28] J. E. Mazo, "Faster-than-Nyquist signaling," *The Bell System Technical Journal*, vol. 54, no. 8, pp. 1451–1462, Oct 1975.
- [29] J. B. Anderson, F. Rusek, and V.   wall, "Faster-than-Nyquist signaling," *Proceedings of the IEEE*, vol. 101, no. 8, pp. 1817–1830, Aug 2013.
- [30] M. Noweir, Q. Zhou, A. Kwan, R. Valivarthi, M. Helaoui, W. Tittel, and F. M. Ghan-nouchi, "Digitally Linearized Radio-Over Fiber Transmitter Architecture for Cloud Radio Access Network's Downlink," *IEEE Transactions on Microwave Theory and Techniques*, vol. 66, no. 7, pp. 3564–3574, 2018.
- [31] M. B. Mello and L. L. Mendes, "Low-Complexity Detection Algorithms Applied to FTMN-GFDM Systems," *IEEE Access*, vol. 10, pp. 101 683–101 696, 2022.
- [32] L. A. M. Pereira, L. L. Mendes, C. J. A. Bastos-Filho, and S. Arismar Cerqueira, "Linearization Schemes for Radio Over Fiber Systems Based on Machine Learning Algorithms," *IEEE Photonics Technology Letters*, vol. 34, no. 5, pp. 279–282, 2022.
- [33] L. A. M. Pereira, L. L. Mendes, C. J. A. Bastos-Filho, and Arismar Cerqueira S. Jr., "Novel Machine Learning Linearization Scheme for 6G A-RoF Systems," *Journal of Lightwave Technology*, 2023.
- [34] L. Pereira, L. L. Mendes, M. A. Cox, and A. Cerqueira, "Open Dataset Initiative for Machine Learning-Based Linearization in Analog Radio over Fiber Systems," *Optics Communications*, p. 131949, 2025.
- [35] K. Katzis, L. Mfupe, and H. M. Hussien, "Opportunities and challenges of bridging the digital divide using 5g enabled high altitude platforms and tvws spectrum," in *2020 IEEE Eighth International Conference on Communications and Networking (ComNet)*. IEEE, 2020, pp. 1–7.
- [36] S. Chakravarty and A. Kumar, "PAPR Reduction of GFDM Signals Using Encoder-Decoder Neural Network (Autoencoder)," *National Academy Science Letters*, pp. 213–217, 2023.
- [37] Y. Rahmatallah and S. Mohan, "Peak-to-average power ratio reduction in OFDM systems: A survey and taxonomy," *IEEE communications surveys and tutorials*, pp. 1567–1592, 2013.
- [38] L. Andrews, R. Phillips, and P. Yu, "Optical scintillations and fade statistics for a satellite-communication system," *applied optics*, vol. 34, no. 33, pp. 7742–7751, 1995.
- [39] Anatel, "Resolu  o n   190," 2021. [Online]. Available: <https://pesquisa.in.gov.br/imprensa/jsp/visualiza/index.jsp?data=06/10/2021&jornal=515&pagina=25&totalArquivos=181>

- [40] R. Menon, R. M. Buehrer, and J. H. Reed, "On the impact of dynamic spectrum sharing techniques on legacy radio systems," *IEEE Transactions on Wireless Communications*, vol. 7, no. 11, pp. 4198–4207, 2008.
- [41] N. Patriciello, S. Lagen, B. Bojović, and L. Giupponi, "NR-U and IEEE 802.11 technologies coexistence in unlicensed mmwave spectrum: Models and evaluation," *IEEE Access*, vol. 8, pp. 71 254–71 271, 2020.
- [42] J. Rinas, R. Seeger, L. Brötje, S. Vogeler, T. Haase, and K.-D. Kammeyer, "A multiple-antenna system for ISM-band transmission," *EURASIP Journal on Advances in Signal Processing*, vol. 2004, pp. 1–13, 2004.
- [43] A. Sahin, R. Yang, E. Bala, M. C. Beluri, and R. L. Olesen, "Flexible DFT-S-OFDM: solutions and challenges," *IEEE Communications Magazine*, vol. 54, no. 11, pp. 106–112, 2016.
- [44] J.-J. Van De Beek, O. Edfors, M. Sandell, S. K. Wilson, and P. O. Borjesson, "On channel estimation in OFDM systems," in *1995 IEEE 45th Vehicular Technology Conference. Countdown to the Wireless Twenty-First Century*, vol. 2. IEEE, 1995, pp. 815–819.
- [45] L. P. Kaelbling, M. L. Littman, and A. W. Moore, "Reinforcement learning: A survey," *Journal of artificial intelligence research*, vol. 4, pp. 237–285, 1996.
- [46] J. Clifton and E. Laber, "Q-learning: Theory and applications," *Annual Review of Statistics and Its Application*, vol. 7, no. 1, pp. 279–301, 2020.
- [47] D. Gündüz, Z. Qin, I. E. Aguerri, H. S. Dhillon, Z. Yang, A. Yener, K. K. Wong, and C.-B. Chae, "Beyond transmitting bits: Context, semantics, and task-oriented communications," *IEEE Journal on Selected Areas in Communications*, vol. 41, no. 1, pp. 5–41, 2022.
- [48] T. Berger, Z. Zhang, and H. Viswanathan, "The CEO problem [multiterminal source coding]," *IEEE Transactions on Information Theory*, vol. 42, pp. 887–902, 1996.
- [49] Y. Oohama, "The rate-distortion function for the quadratic Gaussian CEO problem," *IEEE Transactions on Information Theory*, vol. 44, no. 3, pp. 1057–1070, 1998.
- [50] V. Prabhakaran, D. Tse, and K. Ramachandran, "Rate region of the quadratic Gaussian CEO problem," in *International Symposium on Information Theory, 2004. ISIT 2004. Proceedings*. IEEE, 2004, p. 119.
- [51] D. Slepian and J. K. Wolf, "A coding theorem for multiple access channels with correlated sources," *Bell System Technical Journal*, vol. 52, no. 7, pp. 1037–1076, 1973.
- [52] S.-Y. Tung, *Multiterminal Source Coding*. Cornell University, 1978.
- [53] R. M. Gray, "Toeplitz and circulant matrices: A review," 2006.
- [54] R. Dahlhaus, "On the Kullback-Leibler information divergence of locally stationary processes," *Stochastic processes and their applications*, vol. 62, no. 1, pp. 139–168, 1996.

- [55] E. Hammerich, “Waterfilling theorems for linear time-varying channels and related nonstationary sources,” *IEEE Transactions on Information Theory*, vol. 62, no. 12, pp. 6904–6916, 2016.
- [56] W. Saad, C. Chaccour, C. K. Thomas, and M. Debbah, Eds., *Foundations of Semantic Communication Networks*. Hoboken, NJ, USA: John Wiley & Sons, 2025.
- [57] H. Zhou, Y. Deng, X. Liu, N. Pappas, and A. Nallanathan, “Goal-oriented semantic communications for 6g networks,” *IEEE Internet of Things Magazine*, vol. 7, no. 5, pp. 104–110, 2024.
- [58] C. Shannon and W. Weaver, *The Mathematical Theory of Communication*, ser. Illini books. University of Illinois Press, 1949, no. v. 1. [Online]. Available: <https://books.google.com.br/books?id=OfxtvwEACAAJ>
- [59] W. Yang, H. Du, Z. Q. Liew, W. Y. B. Lim, Z. Xiong, D. Niyato, X. Chi, X. Shen, and C. Miao, “Semantic communications for future internet: Fundamentals, applications, and challenges,” *IEEE Communications Surveys & Tutorials*, vol. 25, no. 1, pp. 213–250, 2023.
- [60] N. Pappas and M. Kountouris, “Goal-oriented communication for real-time tracking in autonomous systems,” in *2021 IEEE International Conference on Autonomous Systems (ICAS)*, 2021, pp. 1–5.
- [61] M. Salimnejad, M. Kountouris, and N. Pappas, “Real-time reconstruction of markov sources and remote actuation over wireless channels,” *IEEE Transactions on Communications*, vol. 72, no. 5, pp. 2701–2715, 2024.
- [62] E. Delfani, G. J. Stamatakis, and N. Pappas, “State-aware timeliness in energy harvesting iot systems monitoring a markovian source,” *IEEE Transactions on Green Communications and Networking*, vol. 9, no. 3, pp. 977–990, 2025.
- [63] A. Zakeri, M. Moltafet, and M. Codreanu, “Semantic-aware real-time tracking of a markov source under sampling and transmission costs,” in *2023 57th Asilomar Conference on Signals, Systems, and Computers*, 2023, pp. 694–698.
- [64] F. Clazzer, A. Munari, G. Liva, F. Lazaro, C. Stefanovic, and P. Popovski, “From 5g to 6g: Has the time for modern random access come?” *arXiv preprint arXiv:1903.03063*, 2019.
- [65] R. D. Yates, Y. Sun, D. R. Brown, S. K. Kaul, E. Modiano, and S. Ulukus, “Age of information: An introduction and survey,” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 5, pp. 1183–1210, 2021.
- [66] I. Kahraman, A. Köse, M. Koca, and E. Anarim, “Age of information in internet of things: A survey,” *IEEE internet of things journal*, vol. 11, no. 6, pp. 9896–9914, 2023.
- [67] P. S. Dester, R. D. Souza, and P. Cardieri, “Analysis of the 1-persistent age threshold slotted aloha,” *IEEE Communications Letters*, vol. 30, pp. 1066–1070, 2026.
- [68] —, “On the deterministic regime of 1-persistent age-threshold slotted aloha in heterogeneous networks,” *IEEE Communications Letters*, pp. 1–1, 2026.

- [69] O. T. Yavascan and E. Uysal, “Analysis of slotted aloha with an age threshold,” *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 5, pp. 1456–1470, 2021.
- [70] Z.-W. Sun, “On disjoint residue classes,” *Discrete mathematics*, vol. 104, no. 3, pp. 321–326, 1992.
- [71] X. Chen, K. Gatsis, H. Hassani, and S. S. Bidokhti, “Age of information in random access channels,” *IEEE Transactions on Information Theory*, vol. 68, no. 10, pp. 6548–6568, 2022.
- [72] Z. Li, W. Wang, C. Wilson, J. Chen, C. Qian, T. Jung, L. Zhang, K. Liu, X. Li, and Y. Liu, “Fbs-radar: Uncovering fake base stations at scale in the wild.” in *NDSS*, 2017.
- [73] L. Arnold, M. Hollick, and J. Classen, “Catch you cause i can: busting rogue base stations using cellguard and the apple cell location database,” in *Proceedings of the 27th International Symposium on Research in Attacks, Intrusions and Defenses*, 2024, pp. 613–629.
- [74] S.-Y. Chang and S. Purification, “Securing cellular availability: The wireless blackhole threat and defense,” in *2025 IEEE Conference on Communications and Network Security (CNS)*. IEEE, 2025, pp. 1–9.
- [75] A. Dabrowski, N. Pianta, T. Klepp, M. Mulazzani, and E. Weippl, “Imsi-catch me if you can: Imsi-catcher-catchers,” in *Proceedings of the 30th annual computer security applications Conference*, 2014, pp. 246–255.
- [76] Z. Zhuang, X. Ji, T. Zhang, J. Zhang, W. Xu, Z. Li, and Y. Liu, “Fbsleuth: Fake base station forensics via radio frequency fingerprinting,” in *Proceedings of the 2018 on Asia Conference on Computer and Communications Security*, 2018, pp. 261–272.
- [77] A. Dabrowski, G. Petzl, and E. R. Weippl, “The messenger shoots back: Network operator based imsi catcher detection,” in *International Symposium on Research in Attacks, Intrusions, and Defenses*. Springer, 2016, pp. 279–302.
- [78] P. K. Nakarmi, M. A. Ersoy, E. U. Soykan, and K. Norrman, “Murat: Multi-rat false base station detector,” *arXiv preprint arXiv:2102.08780*, 2021.
- [79] L. Karaçay, Z. Bilgin, A. B. Gündüz, P. Çomak, E. Tomur, E. U. Soykan, U. Gülen, and F. Karakoç, “A network-based positioning method to locate false base stations,” *IEEE Access*, vol. 9, pp. 111 368–111 382, 2021.
- [80] P. K. Nakarmi, J. Sternby, and I. Ullah, “Applying machine learning on rsrp-based features for false base station detection,” in *Proceedings of the 17th International Conference on Availability, Reliability and Security*, 2022, pp. 1–7.
- [81] S. Purification, K. Park, J. Kim, J. Kim, and S.-Y. Chang, “Wireless link routing to secure against fake base station in 5g,” in *2024 Silicon Valley Cybersecurity Conference (SVCC)*. IEEE, 2024, pp. 1–3.
- [82] S. Purification, S. Wuthier, J. Kim, J. Kim, and S.-Y. Chang, “Fake base station detection and blacklisting,” in *2024 33rd International Conference on Computer Communications and Networks (ICCCN)*. IEEE, 2024, pp. 1–9.

- [83] A. Nassar and Y. Yilmaz, “Reinforcement learning for adaptive resource allocation in fog RAN for IoT with heterogeneous latency requirements,” *IEEE Access*, vol. 7, pp. 128 014–128 025, 2019.
- [84] H. Viswanathan and P. E. Mogensen, “Communications in the 6G era,” *IEEE Access*, vol. 8, pp. 57 063–57 074, 2020.
- [85] M. R. Mahmood *et al.*, “A Comprehensive Rev. on Artificial Intelligence/Machine Learning Algorithms for Empowering the Future IoT Toward 6G Era,” *IEEE Access*, vol. 10, pp. 87 535–87 562, 2022.
- [86] I. F. Akyildiz, A. Kak, and S. Nie, “6G and beyond: The future of wireless communications systems,” *IEEE Access*, vol. 8, pp. 133 995–134 030, 2020.
- [87] Z. Qu, G. Zhang, H. Cao, and J. Xie, “LEO satellite constellation for Internet of things,” *IEEE Access*, vol. 5, pp. 18 391–18 401, 2017.
- [88] H. Jiang, H. Wang, Y. Hu, and J. Wu, “Dynamic user association in scalable ultra-dense leo satellite networks,” *IEEE Transactions on Vehicular Technology*, vol. 71, no. 8, pp. 8891–8905, 2022.
- [89] Lacuna Space. (2023) An ultra-low cost tracking and sensor detection service for small amounts of data. Online; accessed Jul. 26, 2023. [Online]. Available: <https://lacuna.space/>
- [90] Myriota, “Ultralite developer kit,” 2026, acesso em: 29 abr. 2026. [Online]. Available: <https://myriota.com/ultralite-dev-kit/>
- [91] G. M. Capez, S. Henn, J. A. Fraire, and R. Garelo, “Sparse satellite constellation design for global and regional direct-to-satellite IoT services,” *IEEE Transactions on Aerospace and Electronic Systems*, vol. 58, no. 5, pp. 3786–3801, 2022.
- [92] F. A. Tondo, V. D. P. Souto, O. L. Alcaraz López, S. Montejo-Sánchez, S. Céspedes, and R. D. Souza, “Optimal Traffic Load Allocation for Aloha-Based IoT LEO Constellations,” *IEEE Sensors Journal*, vol. 23, no. 3, pp. 3270–3282, 2023.
- [93] V. Fidele Adanvo, S. Mafra, S. Montejo-Sánchez, F. Augusto Tondo, and R. Demo Souza, “Analytical Modeling of Slotted Aloha-Based Direct-to-Satellite-IoT Sensor Networks Over Nakagami-m Fading Channels,” *IEEE Sensors Journal*, vol. 26, no. 2, pp. 3264–3277, 2026.
- [94] F. A. Tondo, M. Afhamisis, S. Montejo-Sánchez, O. L. A. López, M. R. Palattella, and R. D. Souza, “Multiple Channel LoRa-to-LEO Scheduling for Direct-to-Satellite IoT,” *IEEE Access*, vol. 12, pp. 30 627–30 637, 2024.
- [95] F. A. Tondo, S. Montejo Sánchez, and R. D. Souza, “Closeness Centrality-Based Scheduling for IoT Transmissions in LEO Satellite Networks,” in *2025 IEEE Latin Conference on IoT (LCIoT)*, 2025, pp. 335–338.
- [96] R. Ortigueira, S. Montejo-Sánchez, S. Henn, J. A. Fraire, and S. Céspedes, “Satellite visibility prediction for constrained devices in direct-to-satellite iot systems,” *IEEE Sensors Journal*, vol. 24, no. 16, pp. 26 630–26 644, 2024.

- [97] F. A. Tondo, J. M. d. S. Sant’Ana, S. Montejo-Sánchez, O. L. A. López, S. Céspedes, and R. D. Souza, “Nonorthogonal Multiple-Access Strategies for Direct-to-Satellite IoT Networks,” *IEEE Transactions on Aerospace and Electronic Systems*, vol. 61, no. 6, pp. 16 267–16 279, 2025.
- [98] R. R. Vennam, M. Zhang, R. Subbaraman, D. Vashist, and D. Bharadia, “FSMA: Scalable and Reliable LoRa for Non-Terrestrial Networks with Mobile Gateways,” 2025. [Online]. Available: <https://arxiv.org/abs/2510.02800>
- [99] Semtech. (2015) AN1200. 22 LoRa modulation basics. Online.
- [100] T. S. Rappaport, *Wireless Communications: Principles and Practice*. Cambridge University Press, 2024.
- [101] G. Corazza and F. Vatalaro, “A Statistical Model for Land Mobile Satellite Channels and its Application to Nongeostationary Orbit Systems,” *IEEE Transactions on Vehicular Technology*, vol. 43, no. 3, pp. 738–742, 1994.
- [102] J. J. Lopez-Salamanca, L. O. Seman, M. D. Berejuck, and E. A. Bezerra, “Finite-State Markov Chains Channel Model for CubeSats Communication Uplink,” *IEEE Transactions on Aerospace and Electronic Systems*, vol. 56, no. 1, pp. 142–154, 2020.
- [103] S. Herrería-Alonso, M. Rodríguez-Pérez, R. F. Rodríguez-Rubio, and F. Pérez-Fontán, “Improving Uplink Scalability of LoRa-Based Direct-to-Satellite IoT Networks,” *IEEE Internet of Things Journal*, vol. 11, no. 7, pp. 12 526–12 535, 2024.
- [104] R. Ortigueira, S. Montejo-Sánchez, S. Henn, J. A. Fraire, and S. Céspedes, “Satellite Visibility Prediction for Constrained Devices in Direct-to-Satellite IoT Systems,” *IEEE Sensors Journal*, vol. 24, no. 16, pp. 26 630–26 644, 2024.
- [105] D. Croce, M. Gucciardo, S. Mangione, G. Santaromita, and I. Tinnirello, “Impact of lora imperfect orthogonality: Analysis of link-level performance,” *IEEE Communications Letters*, vol. 22, no. 4, pp. 796–799, 2018.
- [106] A. Mahmood *et al.*, “Scalability Analysis of a LoRa Network Under Imperfect Orthogonality,” *IEEE Transactions on Industrial Informatics*, vol. 15, no. 3, pp. 1425–1436, 2019.
- [107] G. Álvarez *et al.*, “Uplink Transmission Policies for LoRa-based Direct-to-Satellite IoT,” *IEEE Access*, vol. 10, pp. 72 687–72 701, 2022.
- [108] S. Herrería-Alonso, M. Rodríguez-Pérez, R. F. Rodríguez-Rubio, and F. Pérez-Fontán, “Improving uplink scalability of lora-based direct-to-satellite iot networks,” *IEEE Internet of Things Journal*, vol. 10, no. 20, pp. C4–C4, 10 2023.
- [109] *LoRa Alliance Technical Committee Regional Parameters Workgroup, Version 1.0.4 and Regional Parameters RP002*, LoRa Alliance, 09 2022. [Online]. Available: <https://resources.lora-alliance.org/technical-specifications/rp002-1-0-4-regional-parameters>
- [110] *LoRa Connect™ 137 MHz to 1020 MHz Long Range Low Power Transceiver*, Semtech Corp, May 2020, accessed: Feb 24, 2026. [Online]. Available: <https://www.semtech.com/products/wireless-rf/lora-connect/sx1276>

- [111] Python Geocoding Toolbox. Accessed: May 01, 2026. [Online]. Available: <https://pypi.org/project/geopy/>
- [112] Elegant Astronomy for Python. Accessed: May 01, 2026. [Online]. Available: <https://pypi.org/project/skyfield/>
- [113] CelesTrack Database. Accessed: May 01, 2026. [Online]. Available: <https://celestrak.org/>
- [114] Lacuna Space, “Lacuna space solutions,” <https://lacuna-space.com/solutions/>, accessed: May, 01, 2026.
- [115] T. L. Marzetta, “Massive MIMO: An introduction,” *Bell Labs Technical Journal*, vol. 20, pp. 11–22, 2015.
- [116] E. Björnson and L. Sanguinetti, “Scalable cell-free massive MIMO systems,” *IEEE Transactions on Communications*, vol. 68, no. 7, pp. 4247–4261, 2020.
- [117] S. Kim, S. Ahn, J. Park, J. Youn, Y. Kwon, and S. Cho, “CPU-cooperative power control scheme for scalable cell-free massive MIMO systems,” *IEEE Transactions on Wireless Communications*, vol. 23, no. 10, pp. 13 904–13 919, 2024.
- [118] F. Li, Q. Sun, X. Ji, and X. Chen, “Scalable cell-free massive MIMO with multiple CPUs,” *Mathematics*, vol. 10, no. 11, p. 1900, 2022.
- [119] E. Björnson and E. Jorswieck, “Optimal resource allocation in coordinated multi-cell systems,” *Foundations and Trends® in Communications and Information Theory*, vol. 9, no. 2–3, pp. 113–381, 2013.
- [120] E. Björnson, J. Hoydis, L. Sanguinetti *et al.*, “Massive MIMO networks: Spectral, energy, and hardware efficiency,” *Foundations and Trends® in Signal Processing*, vol. 11, no. 3–4, pp. 154–655, 2017.
- [121] A. Klautau, P. Batista, N. González-Prelcic, Y. Wang, and R. W. Heath, “5g mimo data for machine learning: Application to beam-selection using deep learning,” in *2018 Information Theory and Applications Workshop (ITA)*. IEEE, 2018, pp. 1–9.
- [122] L. UFPA, “Raymobtime dataset description,” <https://github.com/lasseufpa/Raymobtime/wiki/Raymobtime>, 2021, accessed: August 6, 2025.
- [123] Y. Lu, Y. Wang, and M. Long, “Understanding domain shift by measuring label distribution entropy,” in *NeurIPS*, 2021.
- [124] Y. Chen, T. Lin, and J. Zhang, “Entropy differences as indicators of domain shift in transfer learning,” *IEEE Transactions on Neural Networks and Learning Systems*, vol. 33, no. 9, pp. 4513–4526, 2022.
- [125] H. Kim, H. Lee, T. Kim, and D. Hong, “Cell-free mMIMO systems with dynamic TDD,” in *2022 IEEE 95th Vehicular Technology Conference:(VTC2022-Spring)*. IEEE, 2022, pp. 1–6.
- [126] H. Kim, J. Kim, and D. Hong, “Dynamic tdd systems for 5g and beyond: A survey of cross-link interference mitigation,” *IEEE Communications Surveys & Tutorials*, vol. 22, no. 4, pp. 2315–2348, 2020.

- [127] K. Pedersen, A. Esswie, D. Lei, J. Harrebek, Y. Yuk, S. Selvaganapathy, and H. Helmers, “Advancements in 5g new radio tdd cross link interference mitigation,” *IEEE Wireless Communications*, vol. 28, no. 4, pp. 106–112, 2021.
- [128] K. Ardah, G. Fodor, Y. C. Silva, W. C. Freitas, and F. R. Cavalcanti, “A novel cell reconfiguration technique for dynamic tdd wireless networks,” *IEEE Wireless Communications Letters*, vol. 7, no. 3, pp. 320–323, 2017.
- [129] S. Fukue, H. Iimori, G. T. F. De Abreu, and K. Ishibashi, “Joint access configuration and beamforming for cell-free massive MIMO systems with dynamic TDD,” *IEEE Access*, vol. 10, pp. 40 130–40 149, 2022.
- [130] S. Mashdour, S. Salehi, R. C. de Lamare, A. Schmeink, and J. P. Lima, “Clustering and scheduling with fairness based on information rates for cell-free MIMO networks,” *IEEE Wireless Communications Letters*, vol. 13, no. 7, pp. 1798–1802, 2024.
- [131] S. Mashdour, R. C. de Lamare, and J. P. Lima, “Enhanced subset greedy multiuser scheduling in clustered cell-free massive MIMO systems,” *IEEE Communications Letters*, vol. 27, no. 2, pp. 610–614, 2022.
- [132] F. R. Guimaraes, G. Fodor, W. C. Freitas, and Y. C. Silva, “Pricing-based distributed beamforming for dynamic time division duplexing systems,” *IEEE Transactions on Vehicular Technology*, vol. 67, no. 4, pp. 3145–3157, 2017.
- [133] M. Andersson, T. T. Vu, P. Frenger, and E. G. Larsson, “Joint optimization of switching point and power control in dynamic tdd cell-free massive mimo,” in *2023 57th Asilomar Conference on Signals, Systems, and Computers*, 2023, pp. 988–992.
- [134] J. Rachad, R. Nasri, and L. Decreusefond, “Dynamic-tdd interference tractability approaches and performance analysis in macro-cell and small-cell deployments,” *Annals of Telecommunications*, vol. 75, p. 739–753, 2020.
- [135] S. S. Silva, Y. C. B. Silva, I. M. Guerreiro, V. F. Monteiro, M. D. Carneiro, and A. C. Albuquerque, “User-centric clustering in cell-free networks with dynamic TDD,” in *Proc. Simpósio Brasileiro de Telecomun. e Process. de Sinais (SBrT)*, Oct. 2025, p. 1–5.
- [136] M. D. Carneiro, I. M. Guerreiro, Y. C. B. Silva, S. S. Silva, V. F. Monteiro, and A. C. Albuquerque, “Energy efficiency analysis of power allocation on D-MIMO with dynamic TDD,” in *Proc. Simpósio Brasileiro de Telecomun. e Process. de Sinais (SBrT)*, Oct. 2025, p. 1–5.
- [137] O. T. Demir, E. Björnson, and L. Sanguinetti, “Foundations of user-centric cell-free massive MIMO,” *Foundations and Trends in Signal Processing*, vol. 14, no. 3–4, p. 162–472, 2021.
- [138] P. Baracca, L. G. Giordano, A. Garcia-Rodriguez, G. Geraci, and D. López-Pérez, “Down-link performance of uplink fractional power control in 5G massive MIMO systems,” in *Proc. IEEE Globecom Workshops (GC Wkshps)*, Dec. 2018, pp. 1–6.
- [139] R. Nikbakht and A. Lozano, “Uplink fractional power control for cell-free wireless networks,” in *Proc. IEEE International Conference on Communications (ICC)*, 2019, pp. 1–5.

- [140] C. D'Andrea, A. Garcia-Rodriguez, G. Geraci, L. G. Giordano, and S. Buzzi, "Analysis of UAV communications in cell-free massive MIMO systems," *IEEE Open Journal of the Communications Society*, vol. 1, pp. 133–147, 2020.
- [141] H. D. T. Mathur, "A survey on advanced multiple access techniques for 5g and beyond wireless communications," *Wireless Personal Communications*, pp. 1775–1792, 2021.
- [142] A. Mishra, Y. Mao, O. Dizdar, and B. Clerckx, "Rate-splitting multiple access for 6g—part i: Principles, applications and future works," *IEEE Communications Letters*, vol. 26, no. 10, pp. 2232–2236, 2022.
- [143] D. S. Lakew, A.-T. Tran, A. Masood, N.-N. Dao, and S. Cho, "A review on satellite-terrestrial integrated wireless networks: Challenges and open research issues," in *2023 International Conference on Information Networking (ICOIN)*, 2023, pp. 638–641.
- [144] C. M. Ho, D. S. Lakew, A.-T. Tran, C. Lee, D. T. Hua, and S. Cho, "A review on unmanned aerial vehicle-based networks and satellite-based networks with rsma: Research challenges and future trends," in *2023 International Conference on Artificial Intelligence in Information and Communication (ICAIIIC)*, 2023, pp. 139–142.
- [145] A. Schroeder, M. Roeper, D. Wuebben, B. Matthiesen, P. Popovski, and A. Dekorsy, "A comparison between rsma, sdma, and oma in multibeam leo satellite systems," in *WSA SCC 2023: 26th International ITG Workshop on Smart Antennas and 13th Conference on Systems, Communications, and Coding*, 2023, pp. 1–6.
- [146] A. Schröder, S. Gracla, M. Röper, D. Wübben, C. Bockelmann, and A. Dekorsy, "Flexible robust beamforming for multibeam satellite downlink using reinforcement learning," in *ICC 2024 - IEEE International Conference on Communications*, 2024, pp. 3809–3814.
- [147] B. Clerckx, Y. Mao, R. Schober, and H. V. Poor, "Rate-splitting unifying sdma, oma, noma, and multicasting in miso broadcast channel: A simple two-user rate analysis," *IEEE Wireless Communications Letters*, vol. 9, no. 3, pp. 349–353, 2020.
- [148] C.-L. Hsiao, J.-C. Guey, and R.-J. Chen, "A two-user approximation-based transmit beamforming for physical-layer multicasting in mobile cellular downlink systems," *Journal of the Chinese Institute of Engineers*, vol. 38, no. 6, pp. 742–750, 2015. [Online]. Available: <https://doi.org/10.1080/02533839.2015.1016879>
- [149] Y. Mao, O. Dizdar, B. Clerckx, R. Schober, P. Popovski, and H. V. Poor, "Rate-splitting multiple access: Fundamentals, survey, and future research trends," *IEEE Communications Surveys & Tutorials*, vol. 24, no. 4, pp. 2073–2126, 2023.
- [150] N. Q. Hieu, D. N. Nguyen, D. T. Hoang, and E. Dutkiewicz, "When virtual reality meets rate splitting multiple access: A joint communication and computation approach," *IEEE Journal on Selected Areas in Communications*, vol. 41, no. 5, pp. 1536–1548, 2023.
- [151] T. H. Vu, D. B. da Costa, S. Kim, and Q. V. Pham, "Outage, capacity, and error performance of downlink rsma-based systems: Analysis and resource optimization," *IEEE Transactions on Communications*, vol. 73, no. 8, pp. 6868–6883, 2025.

- [152] T. H. Vu, Q. V. Pham, D. B. da Costa, and S. Kim, "Rate-splitting multiple access-assisted thz-based short-packet communications," *IEEE Wireless Communications Letters*, vol. 12, no. 12, pp. 2218–2222, 2023.
- [153] X. Wang, H. Du, L. Feng, D. Niyato, F. Zhou, Z. Yang, and W. Li, "Resource allocation and user pairing for rate splitting multiple access based wireless networked control systems," *IEEE Transactions on Communications*, vol. 73, no. 7, pp. 4795–4810, 2025.
- [154] Z. Liu, W. Chen, Q. Wu, Z. Li, X. Zhu, Q. Wu, and N. Cheng, "Enhancing robustness and security in isac network design: Leveraging transmissive reconfigurable intelligent surface with rsma," *IEEE Transactions on Communications*, vol. 73, no. 10, pp. 9581–9596, 2025.
- [155] Z. Yang, L. Feng, F. Zhou, K. Xie, X. Qiu, Z. Xiong, C. Yuen, and Z. Han, "User grouping-based hierarchical rate-splitting multiple access for reconfigurable intelligent surface-assisted downlink communication," *IEEE Transactions on Cognitive Communications and Networking*, vol. 11, no. 6, pp. 3811–3826, 2025.
- [156] H. Hashida and B. Di, "Precoding-free hierarchical rate-splitting multiple access via stacked intelligent metasurface," *IEEE Internet of Things Journal*, pp. 1–1, 2026.
- [157] M. Dai, B. Clerckx, D. Gesbert, and G. Caire, "A hierarchical rate splitting strategy for fdd massive mimo under imperfect csit," in *Proceedings of the IEEE 20th International Workshop on Computer Aided Modeling and Design of Communication Links and Networks (CAMAD)*, 2015, pp. 80–84.
- [158] M. Neelam and A. Sundru, "An innovative method for improving the performance of uav-sm-noma for 6g networks over generalized fading channels," *International Journal of Communication Systems*, vol. 38, no. 5, p. e6032, 2025, e6032 IJCS-24-3457.R1. [Online]. Available: <https://onlinelibrary.wiley.com/doi/abs/10.1002/dac.6032>
- [159] M. Ahmed, W. U. Khan, F. N. Al-Wesabi, S. A. Ebad, H. Mesfer Alshahrani, A. Kumar Dutta, B. M. ElHalawany, and X. Li, "Efficient resource management for noma-enabled uav communications in 6g irs-assisted vehicular networks," *IEEE Transactions on Intelligent Transportation Systems*, vol. 27, no. 2, pp. 2706–2715, 2026.
- [160] T.-T. T. Dao, S. Q. Nguyen, H. Nhung-Nguyen, P. X. Nguyen, and Y.-H. Kim, "Performance evaluation of downlink multiple users noma-enable uav-aided communication systems over nakagami-m fading environments," *IEEE Access*, vol. 9, pp. 151 641–151 653, 2021.
- [161] X. Fan, H. Zhou, K. Sun, X. Chen, and N. Wang, "Channel assignment and power allocation utilizing noma in long-distance uav wireless communication," *IEEE Transactions on Vehicular Technology*, vol. 72, no. 10, pp. 12 970–12 982, 2023.
- [162] T. M. Hoang, B. C. Nguyen, H. L. T. Thanh, X. N. Tran, and P. T. Hiep, "Finite block length noma mu pairing uav-enable system: Performance analysis and optimization," *IEEE Transactions on Mobile Computing*, vol. 23, no. 10, pp. 9804–9820, 2024.
- [163] T.-T. Nguyen and A.-V. Nguyen, "Outage analysis of cognitive inspired noma networks in the presence of imperfect sic, cci, and non-id fading channels," *Computer Networks*, vol. 234, p. 109909, 2023. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1389128623003547>

- [164] B. E. Y. Belmekki and M.-S. Alouini, “Noma as the next-generation multiple access in nonterrestrial networks,” *Proceedings of the IEEE*, vol. 112, no. 9, pp. 1303–1345, 2024.
- [165] F. A. C. M. Cardoso, F. R. A. Parente, and M. R. P. da Silva, “Characterization of Power-Domain NOMA Versus MIMO in 6G OFDM Networks,” in *2025 SBMO/IEEE MTT-S International Microwave and Optoelectronics Conference (IMOC)*, Nov. 2025, pp. 660–664.
- [166] B. M. Elhalawany, C. Gamal, A. Elsayed, M. M. Elsherbini, M. M. Fouda, and N. Ali, “Outage analysis of coordinated noma transmission for leo satellite constellations,” *IEEE Open Journal of the Communications Society*, vol. 3, pp. 2195–2202, 2022.
- [167] H. Ayanampudi, T. Shaik, and N. P. Syed, “Performance analysis of noma enabled uav assisted cloud radio access network over composite $\kappa - \mu$ faded channels,” *AEU - International Journal of Electronics and Communications*, vol. 209, p. 156267, 2026. [Online]. Available: <https://www.sciencedirect.com/science/article/pii/S1434841126000737>
- [168] F. R. A. Parente, F. P. Calmon, and J. C. S. S. Filho, “Unified framework for diversity and coding gains over a broad Gaussian class of fading channels,” *IEEE Transactions on Vehicular Technology*, vol. 72, no. 12, pp. 15 916–15 929, 2023.
- [169] B. Lima, N. Fachada, R. Dinis, B. Costa, and M. Boko, “A uav-noma network model under non-ideal conditions of open research software,” *Sensors*, vol. 22, no. 20, p. 7922, 2022.
- [170] F. Rinaldi, H.-L. Maattanen, I. Torsner, P. Pizzi, S. Andreev, A. Iera, Y. Koucheryavy, and G. Araniti, “Non-terrestrial networks in 5g beyond: A survey,” *IEEE Access*, vol. 8, pp. 165 178–165 200, 2020.
- [171] B. Y. Belmekki and M.-S. Alouini, “Noma as the next-generation multiple access in non-terrestrial networks,” in *Proceedings of the IEEE*, vol. 112, no. 10, 2024, pp. 1303–1345.
- [172] A. Han, T. V., and X. Zhang, “Outage performance of noma-based uav-assisted communication with imperfect sic,” in *2019 IEEE Wireless Communications and Networking Conference (WCNC)*, 2019, pp. 1–6.
- [173] P. Bansal, A. Gupta, A. Bansal, I. Khirwar, and P. Garg, “Non-terrestrial network aided noma based underwater acoustic communication system,” in *2025 National Conference on Communications (NCC)*, 2025, pp. 1–6.
- [174] T.-T. Nguyen and A.-V. Nguyen, “Outage analysis of cognitive noma networks in the presence of imperfect sic,” *Physical Communication*, vol. 66, p. 101389, 2025.
- [175] T. T. Dao, S. Q. Nguyen, H. Nhung-Nguyen, P. X. Nguyen, and H. Kim, “Performance evaluation of downlink multiple user mimo-noma with beamforming systems over nakagami- m environments,” *IEEE Access*, vol. 9, pp. 151 641–151 653, 2021.
- [176] M. Neelam and A. Sundru, “An innovative method for improving the performance of uav-sm-noma for 6g networks over cascaded fading channels,” *International Journal of Communication Systems*, vol. 38, no. 5, p. e6032, 2025.

- [177] J. Zhou, C. Yang, J. Huang, S. Yuan, X. Chen, and S. Xie, “Empowered uav networks: Joint clustering and power optimization for energy-efficient federated learning,” *IEEE Networks*, vol. 39, no. 4, pp. 202–2025, 2025.