

UNIVERSIDADE FEDERAL DE GOIÁS
ESCOLA DE ENGENHARIA ELÉTRICA, MECÂNICA E DE
COMPUTAÇÃO

MARCUS VINÍCIUS GONZAGA FERREIRA

**Alocação de Blocos de Recursos em
Redes LTE Utilizando Estimativa de
Limitante de Retardo Através de
Cálculo de Rede**

Goiânia
2015

MARCUS VINÍCIUS GONZAGA FERREIRA

Alocação de Blocos de Recursos em Redes LTE Utilizando Estimativa de Limitante de Retardo Através de Cálculo de Rede

Dissertação apresentada ao Programa de Pós-Graduação da Escola de Engenharia Elétrica, Mecânica e de Computação da Universidade Federal de Goiás, como requisito parcial para obtenção do título de Mestre em Programa de Pós-Graduação em Engenharia Elétrica e de Computação.

Área de concentração: Engenharia de Computação.

Orientador: Prof. Flávio Henrique Teles Vieira, Dr.

Goiânia
2015

Ficha catalográfica elaborada automaticamente
com os dados fornecidos pelo(a) autor(a), sob orientação do Sibi/UFG.

Ferreira, Marcus Vinícius Gonzaga

Alocação de Blocos de Recursos em Redes LTE Utilizando Estimativa
de Limitante de Retardo Através de Cálculo de Rede [manuscrito] /
Marcus Vinícius Gonzaga Ferreira. - 2015.
189 f.: il.

Orientador: Prof. Dr. Flávio Henrique Teles Vieira.

Dissertação (Mestrado) - Universidade Federal de Goiás, Escola de
Engenharia Elétrica (EEEC) , Programa de Pós-Graduação em
Engenharia Elétrica e de Computação, Goiânia, 2015.

Bibliografia. Apêndice.

Inclui siglas, abreviaturas, símbolos, gráfico, tabelas, algoritmos, lista
de figuras, lista de tabelas.

1. LTE. 2. Escalonamento. 3. Retardo. 4. Cálculo de Rede. I. Vieira,
Flávio Henrique Teles, orient. II. Título.

TERMO DE CIÊNCIA E DE AUTORIZAÇÃO PARA DISPONIBILIZAR AS TESES E DISSERTAÇÕES ELETRÔNICAS (TEDE) NA BIBLIOTECA DIGITAL DA UFG

Na qualidade de titular dos direitos de autor, autorizo a Universidade Federal de Goiás (UFG) a disponibilizar, gratuitamente, por meio da Biblioteca Digital de Teses e Dissertações (BDTD/UFG), sem ressarcimento dos direitos autorais, de acordo com a Lei nº 9610/98, o documento conforme permissões assinaladas abaixo, para fins de leitura, impressão e/ou download, a título de divulgação da produção científica brasileira, a partir desta data.

1. Identificação do material bibliográfico: ☒ **Dissertação** ☐ **Tese**

2. Identificação da Tese ou Dissertação

Autor (a):	Marcus Vinícius Gonzaga Ferreira		
E-mail:	marcusviniciusbr@gmail.com		
Seu e-mail pode ser disponibilizado na página?*	☒ Sim	☐ Não	
Vínculo empregatício do autor			
Agência de fomento:		Sigla:	
País:	UF:	CNPJ:	
Título:	Alocação de Blocos de Recursos em Redes LTE Utilizando Estimativa de Limitante de Retardo Através de Cálculo de Rede		
Palavras-chave:	LTE; Escalonamento; Retardo; Cálculo de Rede		
Título em outra língua:	Resource Block Allocation in LTE Using Delay Bound Estimation Through Network Calculus		
Palavras-chave em outra língua:	LTE; Scheduling; Delay; Network Calculus		
Área de concentração:	Engenharia de Computação		
Data defesa: (dd/mm/aaaa)	11/12/2015		
Programa de Pós-Graduação:	Engenharia Elétrica e de Computação		
Orientador (a):	Prof. Dr. Flávio Henrique Teles Vieira		
E-mail:	flavio@emc.ufg.br		
Co-orientador (a):*			
E-mail:			

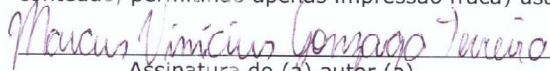
*Necessita do CPF quando não constar no SisPG

3. Informações de acesso ao documento:

Concorda com a liberação total do documento ☒ SIM ☐ NÃO¹

Havendo concordância com a disponibilização eletrônica, torna-se imprescindível o envio do(s) arquivo(s) em formato digital PDF ou DOC da tese ou dissertação.

O sistema da Biblioteca Digital de Teses e Dissertações garante aos autores, que os arquivos contendo eletronicamente as teses e ou dissertações, antes de sua disponibilização, receberão procedimentos de segurança, criptografia (para não permitir cópia e extração de conteúdo, permitindo apenas impressão fraca) usando o padrão do Acrobat.


Assinatura do (a) autor (a)

Data: 13 / 01 / 16

¹ Neste caso o documento será embargado por até um ano a partir da data de defesa. A extensão deste prazo suscita justificativa junto à coordenação do curso. Os dados do documento não serão disponibilizados durante o período de embargo.



MINISTÉRIO DA EDUCAÇÃO
UNIVERSIDADE FEDERAL DE GOIÁS
ESCOLA DE ENGENHARIA ELÉTRICA, MECÂNICA E DE COMPUTAÇÃO
COORDENAÇÃO DE PÓS-GRADUAÇÃO



FOLHA DE APROVAÇÃO

"Alocação de Blocos de Recursos em Redes LTE Utilizando Estimativa de Limitante de Retardo Através de Cálculo de Rede"

MARCUS VINÍCIUS GONZAGA FERREIRA

Dissertação defendida e aprovada pela banca examinadora constituída pelos senhores:

Flávio Henrique Teles Vieira – Orientador (EMC/UFG)

Vinicius da Cunha Martins Borges – INF/UFG

Rodrigo Pinto Lemos – EMC/UFG

Goiânia, 11 de dezembro de 2015

Todos os direitos reservados. É proibida a reprodução total ou parcial do trabalho sem autorização da universidade, do autor e do orientador(a).

Marcus Vinícius Gonzaga Ferreira

Mestrado em andamento em Engenharia Elétrica e de Computação pela Universidade Federal de Goiás - UFG (2015). Pós graduado em Tecnologias para Gestão de Negócios pela Universidade Federal de Goiás - UFG (2012). Graduado em Engenharia de Computação pela Universidade Federal de Goiás - UFG (2010). Técnico em Telecomunicações pelo Centro Federal de Educação Tecnológica de Goiás - CEFET-GO, atual IFG (2006). Serventuário da justiça / Analista de suporte técnico em TI no Tribunal de Justiça do Estado de Goiás desde 2008. Tem experiência na área de Engenharia de Computação, com ênfase em Sistemas de Telecomunicações, atuando principalmente nos seguintes temas: modelagem, predição e controle de tráfego de redes de computadores, redes de computadores, redes sem fio, dispositivos móveis e desenvolvimento de software em linguagem Java.

Dedico este trabalho aos meus familiares e amigos que me incentivaram e ajudaram para que fosse possível sua concretização.

Agradecimentos

Aos meus familiares, pelo incentivo e apoio.

Ao professor Dr. Flávio Henrique Teles Vieira, pela orientação e apoio.

Aos examinadores, pela participação na banca examinadora e sugestões de melhoria deste trabalho.

A todos que contribuíram de alguma forma para a conclusão deste trabalho.

“Dê-me um ponto de apoio e eu moverei a Terra.”

Arquimedes,
Século II a.C..

Resumo

Gonzaga Ferreira, Marcus Vinícius. **Alocação de Blocos de Recursos em Redes LTE Utilizando Estimativa de Limitante de Retardo Através de Cálculo de Rede**. Goiânia, 2015. 189p. Dissertação de Mestrado. Escola de Engenharia Elétrica, Mecânica e de Computação, Universidade Federal de Goiás.

Neste trabalho é proposto um algoritmo de alocação de blocos de recurso para sistemas de comunicação LTE (*Long Term Evolution*) onde é levado em conta o critério de retardo máximo e as restrições impostas pelo esquema MCS (*Modulation and Coding Scheme*) para transmissão *downlink*.

Primeiramente, é proposto um algoritmo de alocação que tenta reduzir o retardo do usuário utilizando a informação do retardo real da rede e a qualidade de transmissão do canal.

Em seguida, é proposto um algoritmo de alocação que considera a qualidade de transmissão do canal e o limitante de retardo, estimado através de Cálculo de Rede utilizando curva de serviço e processo envelope MFBAP (*Multifractal Bounded Arrival Process*), para decidir sobre a alocação de recursos de rádio disponíveis.

São realizadas comparações com outros algoritmos de alocação através de parâmetros de QoS (*Quality of Service*) como retardo médio, vazão total, taxa de perda, índice de justiça (*fairness*) e tempo de processamento, verificando a eficiência do algoritmo proposto.

Palavras-chave

LTE; Escalonamento; Retardo; Cálculo de rede

Abstract

Gonzaga Ferreira, Marcus Vinícius. **Resource Block Allocation in LTE Using Delay Bound Estimation Through Network Calculus**. Goiânia, 2015. 189p. MSc. Dissertation. Escola de Engenharia Elétrica, Mecânica e de Computação, Universidade Federal de Goiás.

In this work we propose an algorithm to allocate resource blocks for LTE (Long Term Evolution) communication systems that takes into account the maximum delay guarantee and MCS (Modulation and Coding Scheme) constraints on the downlink transmission.

At first, we propose an allocation algorithm which tries to reduce user's delay using the information of the network real delay and the channel transmission quality.

Next, we propose an allocation algorithm which considers the channel transmission quality and the delay target, which is estimated through Network Calculus using service curve and MFBAP (Multifractal Bounded Arrival Process) envelope process, in order to decide on the scheduling of available radio resources.

Comparisons with other allocation algorithms are carried out through QoS (Quality of Service) parameters such as average delay, total throughput, loss rate, fairness and processing time, verifying the efficiency of the proposed algorithm.

Keywords

LTE; Scheduling; Delay; Network Calculus

Sumário

Lista de Figuras	12
Lista de Tabelas	19
Lista de Algoritmos	22
Lista de Publicações	23
Lista de Abreviaturas	24
1 Introdução	29
2 Modulação OFDM e LTE	31
2.1 Comunicação Móvel	31
2.2 Modulação OFDM (<i>Orthogonal Frequency-Division Multiplexing</i>)	32
2.2.1 Prefixo Cíclico e Estimativa de Canal	37
2.2.2 Pico de Potência	39
2.2.3 Diagrama de Bloco OFDM	40
2.3 Características Gerais do LTE (<i>Long Term Evolution</i>)	42
2.3.1 Objetivos do LTE	44
2.3.2 Arquitetura da Rede LTE	45
Dispositivos Móveis	45
E-UTRAN, <i>Evolved UMTS Terrestrial Radio Access Network</i>	46
EPC, <i>Evolved Packet Core</i>	48
2.3.3 Arquitetura da Interface de Rádio LTE	49
Canais Lógicos e de Transporte	51
Canal Físico	53
2.3.4 Esquemas de Transmissão LTE	55
<i>Download e Upload</i>	55
Parâmetros de Transmissão LTE	56
Símbolos de Referência e de Sincronização	58
3 Alocação de Recursos em Redes LTE	60
3.1 Alocação de Recursos no <i>Download e Upload</i>	60
3.1.1 Estratégias de Alocação de Recursos no <i>Download</i>	63
3.2 Modelo do Sistema e Formulação do Problema	64
3.3 Algoritmo de Alocação com Garantia de QoS	66
3.3.1 Estimativa de SBs	66
Cálculo da média de ganho de canal de cada usuário	66

	Estimação do número de SBs requeridos por cada usuário	67
3.3.2	Alocação de SBs	67
3.4	Algoritmo de Alocação Baseado em <i>Particle Swarm Optimization</i> (PSO)	69
3.4.1	Otimização PSO Inteira	71
3.4.2	Alocação de Recursos Utilizando PSO	71
3.5	Alocação de Recursos via Modelagem de Tráfego	72
3.5.1	Modelagem Multifractal β MWM Adaptativa	72
3.5.2	Banda Efetiva Utilizando Modelagem β MWM Adaptativa	75
3.5.3	Taxa Mínima Requerida na Alocação de Recursos	76
4	Redução de Retardo na Alocação de Recursos em Redes LTE	80
4.1	Algoritmo de Alocação Proposto	80
4.2	Parâmetros de Simulação	84
4.2.1	Modelagem de Canal	84
4.2.2	Parâmetros do Sistema LTE	86
4.3	Simulações e Resultados	88
4.3.1	Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 1000 TTIs	88
4.3.2	Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 24 Horas de Simulação	98
4.3.3	Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs	102
4.3.4	Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 1000 TTIs	107
4.3.5	Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 24 Horas de Simulação	117
4.3.6	Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs	121
5	Alocação de Recursos em Redes LTE com Estimação de Retardo Utilizando Curva de Serviço e Processo Envelope	126
5.1	Cálculo de Rede Determinístico	126
5.1.1	Processo envelope	128
	Balde Furado Fractal (FLB, <i>Fractal Leaky Bucket</i>)	128
	Processo de Chegada com Limitante Multifractal (MFBAP - <i>Multifractal Bounded Arrival Process</i>)	129
	Análise Comparativa	129
5.1.2	Curva de Serviço	132
5.1.3	Estimativa de Limitante de Retardo	135
5.2	Algoritmo Proposto de Alocação Utilizando Estimação de Retardo	140
5.3	Simulações e Resultados	143
5.3.1	Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 1000 TTIs	143
5.3.2	Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 24 Horas de Simulação	149
5.3.3	Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs	153

5.3.4	Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 1000 TTIs	157
5.3.5	Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 24 Horas de Simulação	163
5.3.6	Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs	167
5.3.7	Controle de Admissão Aplicado em Transmissões em Tempo Real	171
6	Conclusão	176
	Referências Bibliográficas	179
A	Séries Reais de Tráfego TCP/IP - Universidade de Waikato	184

Lista de Figuras

2.1	Medições de tráfego de voz e dados em redes de telefonia móvel em todo o mundo, no período de 2007 a 2013 [40]	31
2.2	Esquema de distribuição de subportadoras OFDM [40]	33
2.3	Esquema de distribuição de subportadoras OFDM [40]	34
2.4	Constelações de diferentes modulações: QPSK (a), 16-QAM (b) e 64-QAM (c) [15]	35
2.5	Espectros sinais OFDM e WCDMA (5MHz) [15]	35
2.6	Tempo de chegada do símbolo em caso de dispersão temporal [40]	37
2.7	Inserção de prefixo cíclico [40]	38
2.8	Diagrama de blocos OFDM	41
2.9	Visão de alto nível da arquitetura LTE	45
2.10	Arquitetura Interna da Rede de Transporte E-UTRAN [40]	48
2.11	Pilha de protocolos da arquitetura de rádio LTE [40]	50
2.12	Mapeamento de Canais de <i>Download</i> [40]	52
2.13	Mapeamento de Canais de <i>Upload</i> [40]	52
2.14	Comparação entre esquemas de transmissão LTE [25]	55
2.15	Estrutura básica de um <i>frame</i> LTE no domínio do tempo e frequência [45]	57
3.1	<i>Channel-dependent scheduling</i> no <i>download</i> para dois usuários, nos domínios do tempo e da frequência [40]	63
3.2	Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 1 até 10	77
3.3	Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 11 até 20	78
3.4	Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 21 até 30	79
4.1	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	90
4.2	Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	93
4.3	Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	94

4.4	Retardo médio entre usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	95
4.5	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	96
4.6	Perda média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	97
4.7	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	100
4.8	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	100
4.9	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	101
4.10	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	101
4.11	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	102
4.12	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	104
4.13	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	105
4.14	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	105
4.15	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	106
4.16	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	106
4.17	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	109
4.18	Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	112

4.19	Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	113
4.20	Retardo médio entre usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	114
4.21	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	115
4.22	Perda média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs	116
4.23	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	119
4.24	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	119
4.25	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	120
4.26	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	120
4.27	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	121
4.28	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	123
4.29	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	123
4.30	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	124
4.31	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	124
4.32	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	125

5.1	A resposta ao impulso da concatenação de dois circuitos lineares é (a) a convolução da resposta ao impulso individual, a curva de formação da concatenação de dois reguladores é (b) a convolução das curvas individuais modeladas [9]	128
5.2	Processo envelope MFBAP, FLB e real para representação dos usuários 1, 5, 10, 15, 20, 25 e 30	130
5.3	Processo envelope MFBAP, FLB e real para representação dos usuários 2, 6, 11, 16, 21 e 26	130
5.4	Processo envelope MFBAP, FLB e real para representação dos usuários 3, 7, 12, 17, 22 e 27	131
5.5	Processo envelope MFBAP, FLB e real para representação dos usuários 4, 8, 13, 18, 23 e 28	131
5.6	Processo envelope MFBAP, FLB e real para representação dos usuários 5, 9, 14, 19, 24 e 29	132
5.7	Curva de serviço do usuário 1 [13]	133
5.8	Curva de serviço por usuário calculada para diferentes cenários em 100 TTIs	135
5.9	Estimativa de limitante de retardo e retardo real médio entre usuários calculado para algoritmo proposto e escalonador Round-Robin, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	137
5.10	Erro Quadrático Médio (EQM) na estimativa de limitante de retardo em função do número de usuários, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	138
5.11	Estimativa de limitante de retardo e retardo real médio entre usuários calculado para algoritmo proposto e escalonador Round-Robin, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	139
5.12	Erro Quadrático Médio (EQM) na estimativa de limitante de retardo em função do número de usuários, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	140
5.13	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	144
5.14	Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	146
5.15	Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	147
5.16	Retardo médio entre usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	147

5.17	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	148
5.18	Perda média do sistema em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	148
5.19	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	150
5.20	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	151
5.21	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	151
5.22	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	152
5.23	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	152
5.24	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	155
5.25	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	155
5.26	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	156
5.27	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	156
5.28	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	157
5.29	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	158
5.30	Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	160
5.31	Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	161

5.32	Retardo médio entre usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	161
5.33	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	162
5.34	Perda média do sistema em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	162
5.35	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	164
5.36	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	165
5.37	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	165
5.38	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	166
5.39	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	166
5.40	Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	168
5.41	Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	169
5.42	Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	169
5.43	Índice de justiça (<i>fairness</i>) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	170
5.44	Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	170
5.45	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	173
5.46	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	173

5.47	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	174
5.48	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	174
5.49	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	175
5.50	Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1.4km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	175
A.1	Série real de tráfego TCP/IP WaikatoVIII-100ms-20110520-000000-0 [48] considerada para usuários 1, 5, 10, 15, 20, 25 e 30	184
A.2	Série real de tráfego TCP/IP WaikatoVIII-100ms-20110623-230233-4 [48] considerada para usuários 2, 6, 11, 16, 21 e 26	185
A.3	Série real de tráfego TCP/IP WaikatoVIII-100ms-20111027-213205-5 [48] considerada para usuários 3, 7, 12, 17, 22 e 27	186
A.4	Série real de tráfego TCP/IP WaikatoVIII-100ms-20110921-000000-0 [48] considerada para usuários 4, 8, 13, 18, 23 e 28	187
A.5	Série real de tráfego TCP/IP WaikatoVIII-100ms-20111029-110001-1 [48] considerada para usuários 5, 9, 14, 19, 24 e 29	188

Lista de Tabelas

2.1	Requisitos de taxa e delay para serviços de dados [40]	42
2.2	Larguras de banda suportadas no LTE [14]	58
3.1	CQI e respectivos esquemas de modulação e taxa de codificação [14]	62
4.1	Parâmetros de modelagem de canal	86
4.2	Valores de critério de retardo máximo adotados	86
4.3	Parâmetros de simulação do sistema LTE	87
4.4	Taxas de bits e SINRs associadas aos MCS [26]	87
4.5	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs	89
4.6	Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs	92
4.7	Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs	98
4.8	Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	99
4.9	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	102
4.10	Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	104
4.11	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	107
4.12	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs	108

4.13	Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator K = 10 e 100) e 1000 TTIs	111
4.14	Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator K = 10 e 100) e 1000 TTIs	117
4.15	Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	118
4.16	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	121
4.17	Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	122
4.18	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	125
5.1	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	144
5.2	Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	146
5.3	Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	149
5.4	Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	150
5.5	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação	153
5.6	Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	154
5.7	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs	157
5.8	Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	158

5.9	Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	160
5.10	Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	163
5.11	Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	164
5.12	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação	167
5.13	Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	168
5.14	Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs	171
5.15	Especificações de retardo recomendadas [10]	171
A.1	Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110520-000000-0 [48]	185
A.2	Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110623-230233-4 [48]	186
A.3	Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20111027-213205-5 [48]	187
A.4	Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110921-000000-0 [48]	188
A.5	Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20111029-110001-1 [48]	189

Lista de Algoritmos

3.1	Algoritmo de escalonamento com garantia de QoS [24]	68
3.2	Algoritmo <i>Particle Swarm Optimization</i> [42]	70
3.3	Algoritmo para Estimativa Adaptativa dos Parâmetros do β MWM [45] [46]	74
4.1	Algoritmo de alocação proposto com critério de retardo	83
5.1	Algoritmo de alocação proposto com estimativa de retardo	142

Lista de Publicações

Abrahão, D.C.; Vieira, F.H.T.; Ferreira, M.V.G., “Algoritmo de Alocação de Blocos de Recursos com Atendimento a Critério de Justiça para Redes sem Fio LTE”. In: XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014, Salvador, BA. Anais do XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014.

Vieira, F.H.T.; Gonçalves, B.H.P.; Ferreira, M.V.G., “Algoritmo Baseado em Enxame de Partículas e Modelagem de Tráfego β MWM para Alocação Dinâmica de Recursos em Redes LTE”. In: XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014, Salvador, BA. Anais do XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014.

Ferreira, M.V.G.; Vieira, F.H.T., “Algoritmo de Alocação de Blocos de Recursos em Redes LTE com Garantia de Retardo Máximo aos Usuários”. 11º Congresso de Pesquisa, Ensino e Extensão (Conpeex). Goiânia. 2014.

Vieira, F.H.T.; Gonçalves, B.H.P.; Rocha, F.G.C.; Lee, L.L.; Ferreira, M.V.G., “Dynamic Resource Allocation in LTE Systems using an Algorithm based on Particle Swarm Optimization and BetaMWM Network Traffic Modeling”. 6th IEEE Latin American Symposium on Circuits and Systems: IEEE LASCAS 2015.

Abrahão, D.C.; Vieira, F.H.T.; Ferreira, M.V.G., “Resource Allocation Algorithm for LTE Networks Using β MWM Modeling and Adaptive Effective Bandwidth Estimation”. In: VI International Workshop on Telecommunications, 2015, Santa Rita do Sapucaí - MG. IEEE IWT, 2015.

Ferreira, M.V.G.; Vieira, F.H.T.; Abrahão, D.C., “Minimizing Delay in Resource Block Allocation Algorithm of LTE Downlink”, VI International Workshop on Telecommunications: IWT 2015.

Ferreira, M.V.G.; Vieira, F.H.T.; Rocha, F.G.C., “Algoritmo Baseado em Estimativa de Retardo Utilizando Curva de Serviço para Alocação Dinâmica de Recursos em Redes LTE”. In: XLVII SBPO 2015 (aceito para publicação), Porto de Galinhas - PE, 2015.

Ferreira, M.V.G.; Vieira, F.H.T.; Castro, M.; Araújo, S.; Rocha, F.G.C., “Alocação de Blocos de Recurso em Redes LTE com Estimativa de Limitante de Retardo através de Cálculo de Rede”. In: XXXIII SBRT 2015, Juiz de Fora - MG. Anais do XXXIII SBRT, 2015.

Ferreira, M.V.G.; Vieira, F.H.T., “Algoritmo de Escalonamento de Recursos para Redes LTE Baseado em Estimativa de Retardo Utilizando Curva de Serviço e Processo Envelope Multifractal”. In: I Workshop em Matemática Industrial, Modelagem e Otimização (WMiMO) 2015, Catalão - GO, 2015.

Lista de Abreviaturas

β MWM: *Beta Multifractal Wavelet Model*
3GPP: *3rd Generation Partnership Project*
ACK: *Acknowledgement*
AMC: *Adaptative Modulation and Coding*
ANR: *Automatic Neighbor Relation*
ARQ: *Automatic Repeat Request*
BCCH: *Broadcast Control Channel*
BCH: *Broadcast Channel*
BER: *Bit Error Ratio*
BLER: *Block Error Rate*
BS: *Base Station*
CCCH: *Common Control Channel*
CDF: *Cumulative Distribution Function*
CDMA: *Code-Division Multiple Access*
CP: *Cyclic Prefix*
CQI: *Channel Quality Indicator*
CRC: *Cyclic Redundancy Check*
DCCH: *Dedicated Control Channel*
DCI: *Downlink Control Information*
DFT: *Discrete Fourier Transform*
DL-SCH: *Downlink Shared Channel*
DRX: *Discontinuous Reception*
DS-CDMA: *Direct-Sequence Code-Division Multiple Access*
DTCH: *Dedicated Traffic Channel*
DVRB: *Distributed Virtual Resource Block*
EDGE: *Enhanced Data Rates For GSM Evolution*
EMM: *Evolved Packet System Mobility Management*
eNodeB: *E-UTRAN Node B, Evolved Node B*
EPA: *Extended Pedestrian A*
EPC: *Evolved Packet Core*

ESM: *Evolved Packet System Session Management*
ETU: *Extended Typical Urban*
E-UTRAN: *Evolved Universal (UMTS) Terrestrial Radio Access Network*
EVA: *Extended Vehicular A*
FDD: *Frequency-Division Duplex*
FDM: *Frequency-Division Multiplex*
FDMA: *Frequency-Division Multiple Access*
FEC: *Forward Error Correction*
FFT: *Fast Fourier Transform*
FLB: *Fractal Leaky Bucket*
GPRS: *General Packet Radio Services*
GSM: *Global System for Mobile Communications*
GTP: *General Packet Radio Service Tunneling Protocol*
HSPA: *High Speed Packet Access*
HSS: *Home Subscriber Server*
IFFT: *Inverse Fast Fourier Transform*
IP: *Internet Protocol*
ISDM: *Integrated Services for Digital Network*
ITU: *International Telecommunication Union*
LB: *Leaky Bucket*
LTE: *Long-Term Evolution*
LVRB: *Localized Virtual Resource Block*
MAC: *Medium Access Control*
MCCH: *Multicast Control Channel*
MCH: *Multicast Channel*
MCS: *Modulation and Coding Scheme*
MFBAP: *Multifractal Bounded Arrival Process*
MIMO: *Multiple-Input Multiple-Output*
MME: *Mobility Management Entity*
MMSE: *Minimum-Mean-Square-Error*
MTCH: *Multicast Traffic Channel*
MWM: *Multifractal Wavelet Model*
NAK: *Negative Acknowledgement*
NAS: *Non-Access Stratum*
NAT: *Network Address Translation*
OFDM: *Orthogonal Frequency-Division Multiplexing*
OFDMA: *Orthogonal Frequency-Division Multiple Access*
PAPR: *Peak-to-Average Power Ratio*

PBCH: *Physical Broadcast Channel*
PCCH: *Paging Control Channel*
PCFICH: *Physical Control Format Indicator Channel*
PCH: *Paging Channel*
PDCCH: *Physical Downlink Control Channel*
PDCP: *Packet Data Convergence Protocol*
PDF: *Probability Density Function*
PDN: *Packet Data Network*
PDSCH: *Physical Downlink Shared Channel*
PDU: *Protocol Data Unit*
PHICH: *Physical Hybrid-ARQ Indicator Channel*
PMCH: *Physical Multicast Channel*
PRACH: *Physical Random Access Channel*
PSO: *Particle Swarm Optimization*
PUCCH: *Physical Uplink Control Channel*
PUSCH: *Physical Uplink Shared Channel*
QAM: *Quadrature Amplitude Modulation*
QoS: *Quality of Service*
QPSK: *Quadrature Phase-Shift Keying*
RAN: *Radio Access Network*
RB: *Resource Block*
RLC: *Radio Link Control*
RNC: *Radio Network Controller*
ROHC: *Robust Header Compression*
RR: *Round Robin*
RRC: *Radio Resource Control*
RS: *Reed-Solomon*
SAE: *System Architecture Evolution*
SB: *Scheduling Block*
SC-FDMA: *Single-Carrier Frequency-Division Multiple Access*
SCTP: *Stream Control Transmission Protocol*
SDU: *Service Data Unit*
S-GW: *Serving Gateway*
SIM: *Subscriber Identity Module*
SINR: *Signal-to-Interference-plus-Noise Ratio*
SIR: *Signal-to-Interference Ratio*
SMS: *Short Message Service*
SNR: *Signal-to-Noise Ratio*

TCP: *Transmission Control Protocol*

TDD: *Time Division Duplex*

TDM: *Time Division Multiplexing*

TSG: *Technical Specification Group*

TTI: *Transmission Time Interval*

UCI: *Uplink Control Information*

UDP: *User Datagram Protocol*

UE: *User Equipment*

UL-SCH: *Uplink Shared Channel*

UMTS: *Universal Mobile Telecommunication System*

USIM: *Universal Subscriber Identity Module*

UTRA: *Universal Terrestrial Radio Access*

VoIP: *Voice over Internet Protocol*

WCDMA: *Wide-Band Code-Division Multiple Access*

Introdução

O aumento do uso da telefonia móvel para o acesso à Internet tem levado a uma crescente demanda por tráfego de dados, tornando necessário a evolução contínua da tecnologia móvel celular. Neste contexto, aparece o sistema LTE (*Long Term Evolution*) como evolução da tecnologia 3G, que tem como objetivo fornecer altas taxas de dados, baixa latência e acesso via rádio otimizado por pacote [2].

O LTE emprega a tecnologia OFDM (*Orthogonal Frequency-Division Multiplexing*) na transmissão de *downlink*, o que permite maior liberdade no escalonamento de canal. O princípio básico do OFDM é a divisão de uma banda larga em múltiplas bandas estreitas utilizadas para a transmissão de informações em paralelo, permitindo o gerenciamento flexível dos recursos de rádio. Um desafio do sistema LTE é suportar um grande número de usuários e satisfazer suas diferentes exigências de taxas de transmissão dado um recurso de rádio limitado [24].

A alocação de recursos de rádio para o sistema LTE tem sido extensamente estudada. Como exemplo, tem-se o algoritmo que utiliza a heurística de otimização PSO (*Particle Swarm Optimization*) [42] a fim de maximizar a vazão total do sistema levando em conta a restrição de taxa mínima de transmissão, e o algoritmo QoS (*Quality of Service*) garantido [24], que objetiva reduzir a complexidade do problema de otimização da vazão total e atender a restrição de taxa mínima de transmissão. O algoritmo PSO é citado na literatura como referência em termos de vazão, porém com alto grau de complexidade computacional e altos valores de retardo. Já o algoritmo QoS garantido apresenta bom desempenho em termos de complexidade computacional e vazão, em detrimento da taxa de perda. Os algoritmos de alocação de recursos de rádio LTE são limitados pela restrição da utilização do MCS (*Modulation Coding Scheme*). Isto é, todos os blocos de escalonamento (SB, *Scheduling Blocks*), unidade básica de alocação, alocados para certo usuário devem usar o mesmo MCS em qualquer intervalo de tempo de transmissão (TTI, *Transmission Time Interval*) com configuração de antena única [42].

Inicialmente, propõe-se um algoritmo que tem como objetivo priorizar o critério de retardo procurando reduzir o retardo (*delay*) médio do sistema. O retardo é um parâmetro de QoS essencial para aplicações em tempo real com taxa de transmissão

variável e requisitos específicos de banda, como serviços de VoIP (*Voice over Internet Protocol*) e de videoconferência, serviços estes que estão cada vez mais em evidência a medida que cresce de forma exponencial o número de dispositivos móveis que utilizam redes LTE. Nesta primeira proposta, é utilizado o valor do retardo observado em rede e a qualidade de transmissão do canal para se efetuar a alocação de recursos, ao contrário do algoritmo QoS garantido e do algoritmo PSO, que consideram unicamente a restrição de taxa mínima de transmissão. O algoritmo proposto é uma variação do algoritmo QoS garantido que mantém sua heurística simples porém com prioridades e critérios diferentes. Sendo assim, procura-se comparar o desempenho em rede do algoritmo proposto com os algoritmos QoS garantido e PSO em termos de parâmetros de QoS como vazão, retardo, índice de justiça e taxa de perda, de maneira a manter desempenho satisfatório em relação aos parâmetros que são pontos fortes nestes algoritmos já bastante referenciados na literatura.

Posteriormente, propõe-se estimar o limitante de retardo de um sistema *downlink* LTE utilizando Cálculo de Rede e conceitos de curva de serviço [13] e processo envelope MFBAP (*Multifractal Bounded Arrival Process*) para o tráfego de redes [44]. Verifica-se o desempenho e a precisão da estimativa de retardo a fim de prover QoS em redes de comunicação, diferentemente da primeira proposta ou mesmo de outros trabalhos da literatura que utilizam valores reais [16] [7] de retardo da rede. Assim, nesta proposta, pode-se atualizar a estimativa de retardo a medida que as características dos dados de tráfego no sistema variam. Será verificado que em muitos quesitos de QoS há melhoria significativa nos índices em relação aos demais algoritmos analisados.

O trabalho está organizado da seguinte forma:

- Capítulo 2: São abordadas as características do sistema de transmissão OFDM e da tecnologia LTE;
- Capítulo 3: São apresentados conceitos de distribuição de recursos em redes LTE;
- Capítulo 4: É proposto o algoritmo de alocação de blocos de recursos em redes LTE com objetivo de reduzir o retardo médio do sistema;
- Capítulo 5: É proposto o algoritmo de alocação de blocos de recursos em redes LTE com objetivo de utilizar estimativa de limitante de retardo a fim de reduzir o retardo;
- Capítulo 6: São apresentadas as considerações finais sobre a dissertação.

Modulação OFDM e LTE

Este capítulo apresenta as principais características da modulação OFDM, descrevendo as vantagens e desvantagens da sua implementação e suas implicações no aperfeiçoamento da tecnologia 4G. Também são abordados os principais aspectos da arquitetura da rede LTE e da sua interface de rádio, ilustrando seus principais componentes, respectivas funções e o modo pelo qual se comunicam, além de apresentar os esquemas adotados na transmissão e multiplexação dos dados entre os elementos da rede e os diversos usuários móveis clientes.

2.1 Comunicação Móvel

Por muito tempo as chamadas de voz dominaram o tráfego nas redes de telecomunicações, porém nos últimos anos o cenário mudou radicalmente. Enquanto a demanda por chamada de voz se manteve estável, o crescimento do tráfego de dados se deu de forma exponencial, como pode ser visto na Figura 2.1.

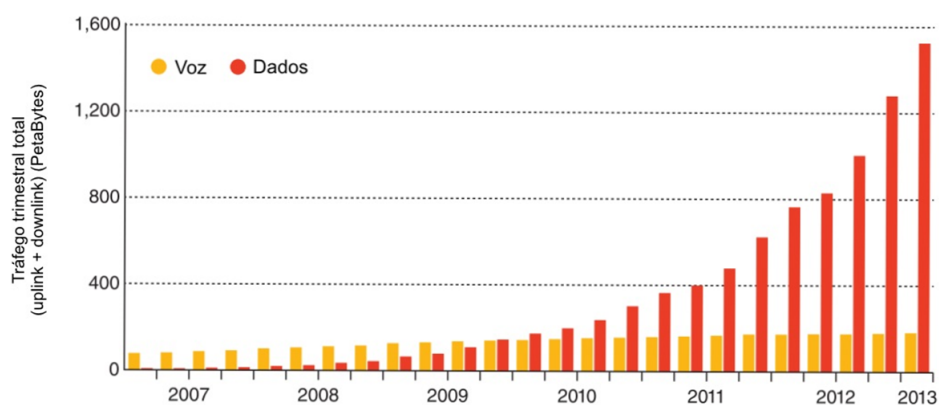


Figura 2.1: *Medições de tráfego de voz e dados em redes de telefonia móvel em todo o mundo, no período de 2007 a 2013 [40]*

Esta tendência deverá continuar, impulsionada pelo aumento da disponibilidade de novas tecnologias de comunicação. Mais importante, porém, foi a introdução de

smartphones com grande capacidade de processamento e praticidade, concebidos para apoiar a criação de aplicativos de terceiros. O resultado foi uma explosão no número e uso de aplicações móveis [43] [40].

O limite teórico da taxa de transmissão C que pode ser alcançada para um sistema de comunicação com largura de banda BW e relação sinal interferência mais ruído $SINR$ (*Signal-to-Interference-plus-Noise Ratio*) é dado pela formula de Shannon [14]:

$$C = B \cdot \log_2(1 + SINR). \quad (2-1)$$

A equação 2-1 mostra que os fatores fundamentais que limitam a taxa de dados são a relação $SINR$ e a largura de banda do canal. Desse ponto de vista, podem-se enxergar várias maneiras de aumentar a capacidade de transmissão de um sistema de comunicação sem fio. Assumindo potência de transmissão constante, um meio de melhorar o $SINR$ é reduzir a distância entre o transmissor e receptor, reduzindo assim a atenuação. Em comunicações móveis isto corresponde à redução no tamanho da célula, criando uma necessidade de mais células para cobrir a mesma área [14] [15] [40].

O aumento da largura de banda acaba por ser ineficiente pois é um recurso caro e escasso. É difícil encontrar largura de banda suficientemente grande hoje, principalmente em baixas frequências, pois estas já se encontram utilizadas por outras tecnologias. Também é importante ressaltar que o uso de grandes larguras de banda tende a aumentar a complexidade e consumo de energia dos equipamentos receptores e transmissores [15] [40].

Talvez a melhor maneira de aumentar as taxas de transmissão seja adoção de uma tecnologia mais robusta, que permita uma aproximação cada vez maior da capacidade do canal teórico, explorando a $SINR$ e a largura de banda disponibilizadas de maneira mais eficiente [14].

2.2 Modulação OFDM (*Orthogonal Frequency-Division Multiplexing*)

OFDM, do inglês *Orthogonal Frequency-Division Multiplexing*, é uma técnica de transmissão sobre múltiplas frequências, baseada na ideia de divisão do canal em várias subportadoras ortogonais entre si com larguras de banda iguais, um tipo especial de FDMA (*Frequency Division Multiple Access*) na qual as subportadoras sobrepõem suas bandas sem causar nenhuma interferência. A ortogonalidade entre as subportadoras torna possível arranjá-las de tal maneira que suas larguras de banda se sobreponham e ainda se consiga recuperar a informação sem interferência entre subportadoras adjacentes [21] [43] [40].

Os esquemas clássicos de transmissão FDMA garantem transmissão sem interferência entre subcanais através da separação destes pela inserção de bandas de guarda entre os mesmos. Essa inserção mantém os subcanais distantes o suficiente para possibilitar posterior separação, mas resulta no uso ineficiente da largura de banda disponível [43] [40].

A modulação OFDM possibilita uma maior eficiência espectral se comparada com técnicas de divisão em frequência convencionais. A Figura 2.2 ilustra bem essa diferença, onde BW é a largura de banda total utilizada para transmissão, N é o número de subportadoras, $SCBW$ é a largura de banda de cada subportadora e R é a taxa de transmissão [43].

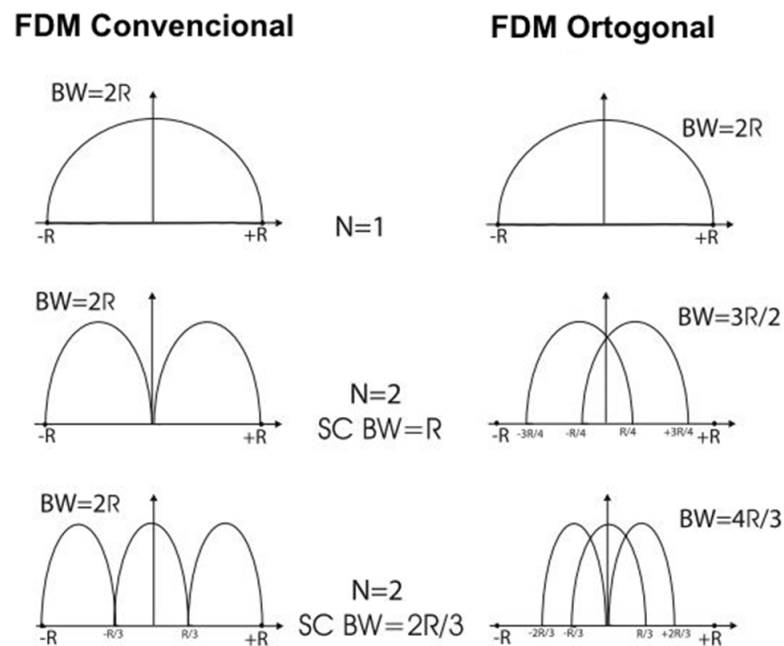


Figura 2.2: Esquema de distribuição de subportadoras OFDM [40]

Entre as principais vantagens do esquema de transmissão OFDM, pode-se citar [40]:

- Simplicidade de implementação: mapeamento de bits pode ser feito através do uso de IFFT/FFT;
- Proteção contra perda de energia;
- Adaptativa no modo que decide a melhor modulação (com mais capacidade) de acordo com o SINR de cada subcanal;
- Redução dos efeitos negativos do canal sobre a transmissão, requerendo receptores mais simples;
- Provê alta velocidade de transmissão com baixo custo;
- Pode suportar acesso dinâmico de pacote;

- Atraente para aplicações de *broadcast*;
- Possibilita fácil integração com antenas inteligentes (sistemas *Multiple-Input Multiple-Output*, MIMO).

OFDM é dita ortogonal devido à ortogonalidade de quaisquer duas subportadoras durante o mesmo intervalo de tempo. Pode ser visto inclusive como uma modulação de um grande número de funções ortogonais. Um bloco de transmissão OFDM é comumente ilustrado como uma grade tempo versus frequência, como mostrado na Figura 2.3 [15].

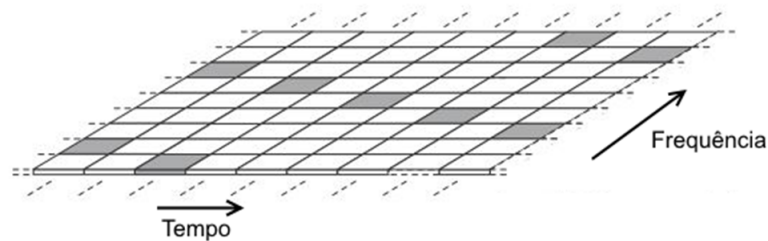


Figura 2.3: Esquema de distribuição de subportadoras OFDM [40]

Ortogonalidade é garantida pela estrutura específica no domínio da frequência na qual o espaço de cada subportadora Δf é exatamente igual à taxa de símbolos $1/T$, sendo T o tempo de um símbolo OFDM. Qualquer corrupção no sinal, por exemplo, pela seletividade em frequência do canal de rádio, pode desfazer essa estrutura em frequência e levar à interferência entre subportadoras [15].

A fim de aumentar a capacidade de transmissão, a OFDM faz o uso de modulações de alta ordem, um meio simples de aumentar o alfabeto de símbolos (mais bits de informação por símbolo). O alfabeto da modulação QPSK (*Quadrature Phase-Shift Keying*) consiste em quatro alternativas diferentes de sinalização, que podem ser representadas por quatro pontos num plano bidimensional, como na Figura 2.4. Nesta figura, observa-se também a constelação das modulações 16-QAM (16 *Quadrature Amplitude Modulation*) e 64-QAM (64 *Quadrature Amplitude Modulation*), com 16 e 64 alternativas diferentes de sinalização, respectivamente [21] [40].

O uso de QPSK permite a transmissão de 2 bits de informação por símbolo, enquanto que 16-QAM e 64-QAM permitem 4 e 6, respectivamente. Sendo assim, consegue-se, em relação à modulação QPSK, um aumento de 2 e 3 vezes no valor máximo de utilização da largura de banda ao se aplicar 16-QAM e 64-QAM em vez de QPSK [15].

Aumentar a ordem de modulação reflete na redução da robustez do sistema a ruídos e interferências, requerendo valores maiores de SINR. O transmissor OFDM decide com base no valor de SINR de cada subportadora qual modulação adotar, visando sempre a maior taxa de transmissão. Valores de SINR baixos requerem modulações mais robustas, enquanto valores de SINR mais altos permitem modulações com taxas mais altas de transmissão [15].

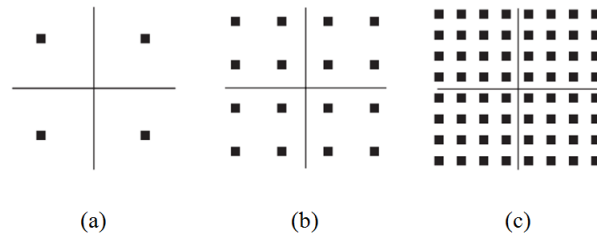


Figura 2.4: Constelações de diferentes modulações: QPSK (a), 16-QAM (b) e 64-QAM (c) [15]

A quantidade de subportadoras varia bastante, e depende da largura de banda disponível e do espaço de cada subportadora. Este é determinado de acordo com o ambiente em que o sistema será aplicado. Aspectos como seletividade de frequência e efeito Doppler determinam o espaço mínimo para a subportadora. Sendo N o número de subportadoras e Δf a largura de banda de uma única subportadora, uma vez determinado Δf , o cálculo do número de subportadoras se dá sabendo que a relação $N \cdot \Delta f$ é igual à largura de banda total do canal [15] [40].

Deve-se levar em conta que o espectro do sinal básico OFDM, mostrado na Figura 2.5, decai muito vagarosamente, se comparado, por exemplo, ao sinal WCDMA (Wide-Band Code-Division Multiple Access), que utiliza como método de múltiplo acesso o CDMA de Sequência Direta (DS-CDMA), com os vários terminais compartilhando uma mesma banda de frequências, porém utilizando códigos diferentes de espalhamento espectral [15]. Essa emissão fora de banda é causada pelo formato do sinal como um pulso retangular. Na prática, um simples filtro pode ser usado para suprimir essa emissão fora de banda, consumindo em torno de 10% da largura de banda total em forma de banda de guarda. Por exemplo, para uma alocação de 10 MHz, o produto $N \cdot \Delta f$ deve ser igual a 9 MHz. Sistemas LTE utilizam subportadoras de 15 kHz, para essa largura de banda apresentada como exemplo, o número de subportadoras fica na ordem de 600 [15] [40].

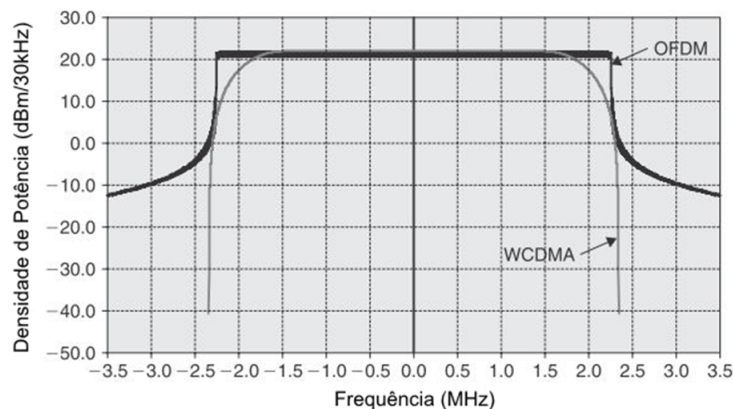


Figura 2.5: Espectros sinais OFDM e WCDMA (5MHz) [15]

Umas das chaves do sistema OFDM é o uso da Transformada Discreta Inversa de

Fourier no transmissor (IDFT) e da Transformada Discreta de Fourier no receptor (DFT). Sendo T o tempo de duração de um símbolo, é devido à sua característica de $\Delta f = 1/T$ que a modulação e demodulação podem ser implementadas com baixa complexidade computacional por essas transformadas [8] [40].

Para transmissão, o sistema OFDM mapeia as mensagens em uma sequência de símbolos PSK ou QAM, subsequentemente convertendo-os em N fluxos paralelos. Cada um dos N símbolos da conversão em série para paralelo é feita para uma subportadora diferente. Cada DFT e IDFT requer N^2 operações, se cada símbolo de saída fosse calculado separadamente. Realizando esse cálculo simultaneamente e tomando vantagem das propriedades cíclicas dos multiplicadores $e^{\pm j2\pi nk/N}$, a Transformada Rápida de Fourier (FFT) reduz a complexidade operacional para a ordem de $N \cdot \log N$ [8] [15].

Existem vários algoritmos de FFT com diferentes ordens de entrada e saída e diferentes usos de memória temporária. Elas podem ser implementadas por microprocessadores especializados em processamento digital para fins gerais ou em circuitos especiais, dependendo principalmente da taxa de informação a ser alcançada [8].

O processamento FFT é mais eficiente quando N é uma potência de dois. Caso não seja, é comum adicionar símbolos de valor zero para o bloco de entrada de modo a que a vantagem de utilizar tal comprimento do bloco na FFT seja alcançada. A quantidade de N , que determina a taxa de símbolos de saída é, então, muitas vezes maior que o número de subportadoras [8] [40].

Os filtros analógicos na saída do transmissor e de entrada do receptor devem limitar em banda os respectivos sinais para $1/T$, onde T é o tempo de duração de um símbolo OFDM. Filtros passa-baixa podem ser utilizados em vez dos filtros passa-banda. O filtro de transmissão elimina emissão fora de banda que poderia interferir com outros sinais. O filtro no receptor é essencial para evitar efeitos de *aliasing*. O sinal transmitido tem um espectro contínuo cujas amostras em frequências são espaçadas de $1/T$. Em particular, o espectro de cada subportadora é da forma $\text{sinc}(1/T)$, cujo valor central é a entrada dos valores dos símbolos, e valores nulos ocorrem exatamente nos valores das frequências centrais das outras subportadoras [8].

As operações de recepção são essencialmente opostas às que ocorrem no transmissor. Um conjunto crítico de operações a serem realizadas são as de sincronização de frequência da portadora, a taxa de amostragem e bloco de enquadramento. Existe um atraso mínimo de $2 \cdot T$ através do sistema por causa das funções de montagem dos blocos de dados no transmissor e receptor [8] [40].

2.2.1 Prefixo Cíclico e Estimativa de Canal

Mesmo sendo cada subportadora estreita no nível de se desconsiderar, dentro de cada uma, a seletividade de frequência dos canais de rádio, este efeito pode causar problemas quanto à manutenção da ortogonalidade entre as subportadoras. Diferentes subportadoras que sofrem com essa dispersão temporal chegarão ao receptor com atrasos diferentes e a ortogonalidade entre elas será afetada. O intervalo de um símbolo irá sobrepor-se com o limite de outro de determinada subportadora que percorreu um caminho diferente para chegar ao receptor, como ilustrado na Figura 2.6 [15] [40].

Em ambientes com essa característica haverá interferência entre símbolos e entre subportadoras. Mesmo que somente uma subportadora seja afetada, por consequência, sua relação com as demais também será, levando à perda de ortogonalidade (interferência entre subportadoras). Devido ao estreito espaço entre as subportadoras, uma transmissão OFDM terá grande sensibilidade a esse tipo de interferência entre subportadoras[11] [15] [40].

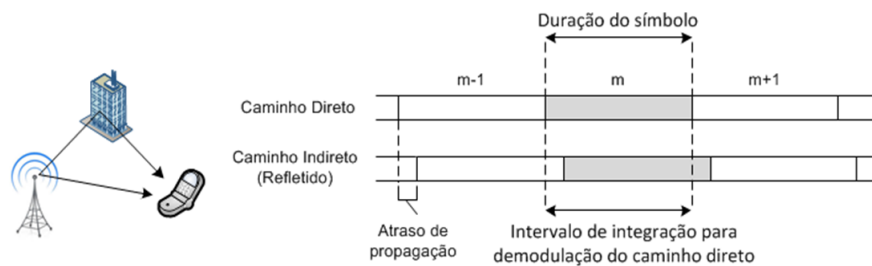


Figura 2.6: Tempo de chegada do símbolo em caso de dispersão temporal [40]

Para lidar com esse problema, a inserção de um intervalo de guarda, chamado prefixo cíclico, é utilizada nas transmissões OFDM, como mostrado na Figura 2.7. Essa técnica consiste em copiar a última parte do símbolo OFDM e inserir no seu início, incrementando assim o tamanho do símbolo. Isso garante que as réplicas atrasadas de um sinal OFDM sempre formarão um símbolo completo dentro da janela FFT. A inserção é feita na saída temporal discreta do transmissor IFFT [15] [43] [40].

Na transformada de Fourier, todos os componentes resultantes do sinal original são ortogonais entre si. Assim, fornecendo a periodicidade do sinal de fonte de OFDM, a inserção de prefixo cíclico assegura que as subportadoras subsequentes continuem ortogonais. No receptor, os dados duplicados são descartados antes de qualquer processamento. Desde que a extensão da dispersão temporal seja menor que o tamanho do prefixo inserido, todas as reflexões de símbolos anteriores são removidas e a ortogonalidade é garantida [15] [43] [40].

As desvantagens da aplicação da técnica de inserção de prefixo cíclico se resumem à perda de fração da potência do sinal utilizada na transmissão, pois parte da

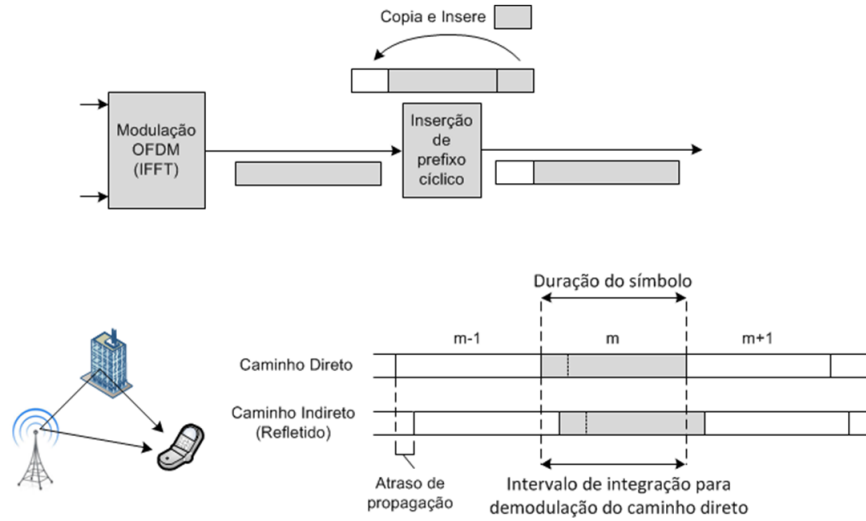


Figura 2.7: *Inserção de prefixo cíclico* [40]

potência total é utilizada na transmissão do prefixo, e à perda na largura de banda, pois reduz a taxa total de transmissão sem redução da largura de banda [15]. Sendo T a duração de um símbolo OFDM e T_{PC} o tamanho do intervalo do prefixo cíclico adotado, a perda $SNR_{perda_{PC}}$ na relação sinal ruído devido à inserção de prefixo cíclico pode ser medida através da seguinte equação:

$$SNR_{perda_{PC}} = -10 \cdot \log_{10} \left(1 - \frac{T_{PC}}{T} \right). \quad (2-2)$$

O tamanho T_{PC} do prefixo cíclico deve cobrir o máximo tamanho da dispersão temporal esperada no canal. Em contraposição, o aumento de T_{PC} sem proporcional redução de Δf implica em sobrecarga adicional tanto na potência quanto na largura de banda requerida. Em termos de perda de potência, com o crescimento da célula o sistema celular torna-se mais limitado, pois a dispersão temporal máxima também cresce junto [43] [40].

Em geral, há uma compensação na perda de potência devido à inserção de prefixo cíclico e à corrupção do sinal causada pela dispersão temporal residual. Em determinado ponto, o aumento no tamanho do prefixo para reduzir a corrupção do sinal não se justifica, pois causa uma perda de potência maior que pela própria corrupção do sinal. Logo, a inserção de prefixo cíclico não deve necessariamente cobrir todo o intervalo de dispersão temporal de determinado canal [40] [15].

Para lidar com diferentes ambientes com diferentes características de célula e potência, alguns sistemas OFDM suportam múltiplos tamanhos de prefixos cíclicos. Prefixos menores utilizados em células pequenas minimizam a sobrecarga gerada, e prefixos de maior duração são aplicados para cobrir uma eventual maior dispersão temporal máxima [15] [40].

Assumindo um prefixo cíclico suficientemente largo, a convolução linear do sinal de entrada com a resposta ao impulso de um canal de rádio que sofre dispersão temporal pode ser vista como uma convolução circular durante o intervalo de demodulação T . Para recuperar um símbolo, é necessário retirar o prefixo cíclico inserido e multiplicar o sinal restante pelo complexo conjugado H_k^* , a Transformada de Fourier da resposta ao impulso do canal da subportadora k . Para isso, é necessário estimar as respostas ao impulso $h(t)$ [15] [40].

Na modulação OFDM, o transmissor modula a sequência de bits de mensagem em símbolos PSKQAM, realiza a IFFT nos símbolos para convertê-los em sinais no domínio do tempo, e os transmite através de um canal de comunicação sem fio. O sinal recebido é usualmente distorcido pelas características do canal. Para recuperar os bits transmitidos, os efeitos do canal devem ser estimados e compensados no receptor [11].

A estimativa do canal é feita diretamente através do uso de símbolos pilotos. Os símbolos já conhecidos pelo receptor são enviados em intervalos regulares na grade OFDM. Usando o conhecimento destes símbolos o receptor pode estimar o canal envolto do símbolo piloto. Logo, estes símbolos devem ter densidade suficiente em tempo e frequência para que sejam capazes de prover uma boa estimativa do canal, mesmo em casos de canais sujeitos a altos níveis de seletividade temporal e/ou em frequência. Vários algoritmos podem ser utilizados para esse fim, desde uma simples interpolação linear a modelos mais complexos, como a estimação *Minimum-Mean-Square-Error* (MMSE) [40] [15].

2.2.2 Pico de Potência

Um dos problemas de transmissão via multiportadoras é a variação na potência instantânea de transmissão, que implica numa redução de eficiência e alto consumo de energia nos terminais. Alternativamente, a saída do amplificador tem de ser reduzida (intervalo de operação reduzido), degradando seu desempenho, além de dificultar a implementação dos conversores digital para analógico e analógico para digital [15] [40] [43].

A relação pico-média de potência (*Peak-to-Average Power Ratio*, PAPR) é definida pela razão entre o pico de potência (potência da onda de maior amplitude) e a potência média do sinal. Essa relação cresce junto com o número de subportadoras. Quando N sinais são transmitidos com a mesma fase, eles produzem um pico de potência igual a N vezes a potência média destes sinais [40] [35].

No intuito de reduzir o problema, várias técnicas são propostas. Uma das mais simples é cortar o sinal, limitando a amplitude a um valor máximo. Embora simples, acarreta alguns problemas que distorcem a amplitude do sinal, produzindo um tipo de

auto interferência que degrada a taxa de erro de bits (BER, *Bit Error Ratio*) e aumenta o nível de emissão fora de banda [35].

Outras soluções para o problema de pico de potência se dão baseadas em codificação. Pode-se adotar uma codificação que só produza símbolos OFDM abaixo de certo nível de PAPR. Quanto menor o PAPR desejado, menor a taxa de bits alcançada. A técnica *scrambling* reduz uma PAPR também aplicando codificação, mas de uma maneira diferente. Dado um símbolo OFDM, a sequência de entrada é embaralhada em várias sequências diferentes e aquela que produz o menor valor de PAPR é transmitida. Depois da demodulação OFDM do lado do receptor o desembaralhamento e subsequente decodificação são realizados para todas as possíveis sequências, sendo que somente uma das decodificações feitas proverá um resultado correto. Essa técnica não reduz a PAPR a um limite fixo, mas diminui a probabilidade da ocorrência de altos níveis de PAPR [15] [40] [35].

2.2.3 Diagrama de Bloco OFDM

Com base no diagrama Figura 2.8, destacam-se as seguintes atividades importantes presentes no modelo de transmissão OFDM [40]:

- A qualidade de um canal varia no domínio da frequência, logo subportadoras num canal OFDM podem sofrer de profundo desvanecimento, com SINR abaixo do requerido, apresentando assim altas taxas de erros e diminuindo a qualidade geral da transmissão. A pré-correção de erros pode evitar uma redução desnecessária da SINR e consequentemente na taxa de transmissão. Vários são os códigos capazes de implementar essa pré-correção, como RS (*Reed-Solomon*), código convolucional, *turbo code*, entre outros [11] [15].
- A pré-correção tem limite de capacidade, um número máximo de erros corrigíveis garantidos por *codeword*, e pode não ser efetiva para erros causados pelo desvanecimento intenso de algumas subportadoras. Na prática, o *interleaving* é implementado para tornar esses erros aleatórios. Distribuindo esses bits em diferentes frequências, cada bit de informação estará sujeito a uma mesma qualidade de transmissão. Esse processo é aplicado após a codificação e é essencial na transmissão OFDM, que acaba se beneficiando da diversidade de frequência em um canal seletivo à frequência [11] [15].
- IFFT e FFT são usadas respectivamente para modulação e demodulação das constelações de símbolos nas subportadoras. Na entrada do processamento IFFT, N pontos da constelação de dados estão presentes, sendo N o número de pontos DFT. As constelações variam de acordo com a modulação utilizada na subportadora (PSK ou QAM). As N saídas do IFFT formam a banda base, carregando os símbolos (dados)

no grupo de N subportadoras. Por causa de alguns filtros passa-baixa requeridos na conversão analógico-digital (e vice-versa), nem todas subportadoras são utilizadas para dados. As que se encontram perto da frequência de Nyquist serão atenuadas e distorcidas por esses filtros, tornando-se impróprias para transmissão [35].

- A sincronização é um problema chave para o projeto de um receptor OFDM robusto. Tanto no tempo quanto na frequência, ela é importante para identificação do início do símbolo e do alinhamento dos osciladores do modulador e demodulador. A modulação OFDM é altamente sensível a erros de sincronização, especialmente na frequência. Os osciladores de frequência do transmissor e receptor devem agir o mais igualmente possível, e pequenos erros levam à rotação de fase dos pontos da constelação de dados. Quanto maior o número de portadoras, maior a precisão requisitada entre os osciladores. Se a sincronização não tem precisão suficiente, a ortogonalidade é parcialmente perdida [11].

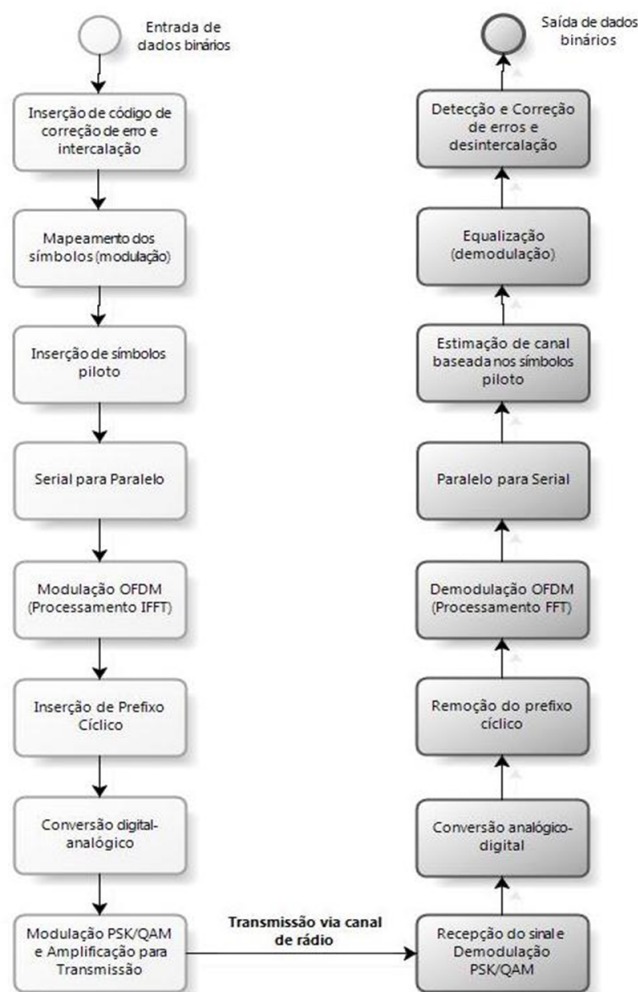


Figura 2.8: Diagrama de blocos OFDM

2.3 Características Gerais do LTE (*Long Term Evolution*)

Os recentes avanços tecnológicos possibilitam o processamento rápido em dispositivos móveis, tornando-os multiuso e capazes de oferecer uma grande gama de serviços processando grandes quantidades de informação em pouco tempo. Os processadores digitais modernos possibilitam novos e mais poderosos sistemas provendo novos serviços, além de conseguirem a melhoria da qualidade e redução do custo nos serviços já existentes. Pode-se citar, por exemplo, serviços que requerem grande quantidade de dados, como aplicações para *download* de arquivos, e outros que requerem baixo *delay*, como televisão digital e voz. Os requisitos dos serviços mais comuns podem ser observados na Tabela 2.1. Deve-se levar em conta também que há inúmeras possibilidades de serviços que podem surgir e que as tecnologias devem ser capazes de suprir suas necessidades [40] [15].

Serviço	Taxa de transferência	Delay
Jogos em tempo real	Baixa	Baixo
Voz	Baixa	Baixo
<i>Download/Upload</i> interativo	Alta	Baixo
<i>Download/Upload</i> em segundo plano	Alta	Alto
Televisão	Média (<i>Streaming Downlink</i>)	Baixo

Tabela 2.1: Requisitos de taxa e delay para serviços de dados [40]

As tecnologias de comunicação móvel são desenvolvidas para atender essas demandas de comunicação sem fio. A Internet e seus serviços baseados no protocolo IP (*Internet Protocol*) estão ocupando cada vez mais espaço no ambiente sem fio, essencialmente por causa da versatilidade do IP, que independe do contexto chamado. Isso pode ser observado na substituição de serviços de comutação de circuitos bem sucedidos, como os serviços de voz, por exemplo, para serviços baseados em pacotes IP, como VoIP (*Voice over IP*, voz sobre IP) [40] [15].

No intuito de suprir as necessidades crescentes de aumento da capacidade de transmissão da rede *wireless* e sabendo das limitações do esquema UMTS (*Universal Mobile Telecommunication System*), a 3GPP (*3rd Generation Partnership Project*) lançou em 2004 a ideia de um novo padrão para sistemas de transmissão via ondas de rádios, o LTE. Do inglês *Long Term Evolution*, ou Evolução de Longo Prazo, essa tecnologia representa a transição da terceira para a quarta geração dos sistemas celulares. Tem como objetivo o aumento da capacidade da rede, com foco na minimização da complexidade dos equipamentos da rede e do usuário, a fim de permitir implantação de espectro flexível nas bandas existentes e futuras, além de habilitar coexistência com outras tecnologias 3GPP de acesso via rádio. Diferente dos anteriores, a tecnologia LTE desenvolveu-se para poder

operar em novos e mais complexos arranjos do espectro, com a possibilidade dos novos projetos não precisarem atender terminais legados [15] [39] [43].

A tecnologia 3G é baseada no sistema UMTS e suas atualizações suportam tanto operações de comutação de circuitos, como de transmissão de pacotes. O novo modelo dá suporte somente a transmissão de informação baseada em pacotes, com objetivo de fornecer altas taxas de transferência, baixa latência e suporte a tecnologias *wireless* emergentes. A única exceção é o SMS (*Short Message Service*), que é transportado por mensagens de sinalização. Essas características trazem uma enorme simplificação no desenho e implementação da arquitetura, além de acompanhar a tendência mundial de migração para transmissão via protocolo IP [40] [15] [39] [43].

A interface de rádio LTE foi construída de forma a suprimir os efeitos de desvanecimento multipercurso na transmissão do sinal através da transmissão via múltiplas portadoras. Toda sua rede é otimizada para tráfego IP, sem suporte para ISDN (*Integrated Services for Digital Network*), isto é, sem requisitos para suporte a serviços de comutação de circuitos GSM (*Global System for Mobile Communications*). Sua ideia é aproveitar o melhor de WCDMA (*Wideband Code Division Multiple Access*) e HSPA (*High Speed Packet Access*) e refazer as partes que devem ser atualizadas em função das novas exigências [39].

Por utilizar a tecnologia OFDM no *download*, o LTE pode adaptar-se a diferentes disponibilidades de espectro ajustando o número de subportadoras à largura de banda disponível, sem alterar os parâmetros próprios das subportadoras. Os dispositivos móveis compatíveis devem suportar todas as possibilidades de largura de banda, que pode variar entre 1.25 e 20 MHz [40] [39].

Um fator que foi levado muito em conta no projeto LTE foi a simplificação. A transmissão e as interfaces entre os nós da rede são baseadas em IP, incluindo a conexão entre as estações base, de forma a tornar a estrutura física completamente transparente e permitir que vários outros protocolos possam atuar nas camadas inferiores [39].

Definiu-se um menor número de componentes na rede lógica e física, que juntamente com procedimentos de sinalização otimizados para o estabelecimento da conexão, e outros de interface e gestão da mobilidade, resultaram em uma redução significativa do tempo de atraso, melhorando ainda mais a experiência do usuário. O tempo necessário para se conectar à rede é da ordem de apenas algumas centenas de milissegundos e a entrada nos estados de economia de energia pode ser alcançada rapidamente [39].

O desenvolvimento da interface está estritamente acoplado ao projeto SAE (*System Architecture Evolution*), a fim de definir a arquitetura do sistema e núcleo da rede. A capacidade de taxa dos equipamentos de usuário compatíveis com LTE é consideravelmente maior que em tecnologias 3GPP anteriores, o que possibilita altos

picos de transferência no *downlink* e *uplink* e flexibilidade aos provedores para adaptarem seus serviços de acordo com o espectro disponível, ou mesmo a habilidade de iniciar um serviço limitado de baixo custo, com possibilidade de aumento de capacidade de acordo com a disponibilidade de espectro [40] [43].

2.3.1 Objetivos do LTE

A 3GPP especificou requisitos bastante exigentes para LTE, conquistados através de melhorias na camada física, como o uso de OFDM e MIMO, e simplificação de protocolos e funcionalidades. A nova interface de rádio, chamada E-UTRAN (*Evolved Universal Terrestrial Radio Access Network*), melhora substancialmente a vazão do usuário final e a capacidade da célula, além de reduzir a latência, trazendo melhor experiência ao usuário. Sua arquitetura de núcleo suporta a E-UTRAN via redução do número de elementos de rede e simplificação de suas funcionalidades, tendo sempre em foco permitir operações de mobilidade transparentes ao usuário final [40] [43].

Segundo [43], dentre os principais requisitos da tecnologia LTE, pode-se citar:

- Picos de transferência para *download* de 100Mb/s (5 bps/Hz) e para *upload* de 50Mb/s (2.5 bps/Hz), ambos para 20MHz de largura de banda alocada;
- Otimização para baixa mobilidade (velocidades até 15km/h). Suporte de alta performance para alta mobilidade (entre 15 e 120 km/h). Entre 120 e 350 km/h a mobilidade deve também ser mantida, mesmo acima de 500 km/h, dependendo da banda de frequência utilizada. O suporte para alta mobilidade ocasionou menores intervalos de transmissão e mecanismos de retorno mais rápidos se comparados à UMTS;
- A operação de *handover* deve ser de tal maneira que o tempo de interrupção seja menor que quando realizada via comutação de circuitos 2G, e *handovers* para sistemas 2G/3G devem ocorrer de forma transparente para o usuário (*seamless*);
- Baixa latência no plano de controle: tempos de transição de status do dispositivo móvel reduzidos;
- Baixa latência no plano de usuário: menos de 5 ms para usuário único com único *stream* de dados;
- Vazão média do usuário/MHz e eficiência de espectro em rede carregada três ou quatro vezes maior que *Release 6* no *download*, e duas ou três vezes maior que *Release 6* no *upload*;
- Vazão, eficiência de espectro e mobilidade devem ser garantidos para células de 5 km, com leve degradação para células de 30 km;
- Redução da complexidade e consumo de energia dos equipamentos dos usuários e da rede;
- Migração de baixo custo UMTS para LTE;

- Serviços *multicast* aprimorados;
- Suporte a qualidade de serviço (*Quality of Service*, QoS) fim-a-fim aprimorado;
- Redução de redundâncias na arquitetura.

Diante de tais aspectos é importante ressaltar que a taxa de dados alcançada na prática é muito menor que a teórica proposta, pois ela depende de vários fatores como o total de espectro e a potência de transmissão utilizados pela estação base, além da distância entre a torre e o dispositivo móvel e a interferência entre estações vizinhas [39].

Muitas das técnicas adotadas na tecnologia LTE já tinham sido introduzidos na evolução UMTS, aplicados a fim de satisfazer a crescente demanda por altas velocidades de transmissão. Pode-se citar MIMO (*Multiple Input, Multiple Output*), especificado para HSPA+ (*Evolved High-Speed Packet Access*) no *Release 7*, a solução *One-Tunnel* para redução do número de nós na rede no *stream* de dados do usuário e o uso de IP em todas interfaces. Esse fator possibilita a fácil integração com as tecnologias predecessoras, tornando possível inclusive a combinação de estações base para suporte a diferentes tecnologias [40] [39].

2.3.2 Arquitetura da Rede LTE

A tecnologia LTE, desenvolvida pela 3GPP, segue o mesmo padrão de arquitetura GSM/UMTS, porém com um número reduzido de elementos. Uma visão geral de alto nível da arquitetura pode ser visualizada na Figura 2.9. Ela é composta pelos dispositivos móveis dos usuários, a rede de rádio (*Evolved UMTS Terrestrial Radio Access Network*, E-UTRAN), e o núcleo da rede (*Evolved Packet Core*, EPC). As interfaces entre as diferentes partes são denotadas por Uu, S1 e SGi, respectivamente [40] [14].

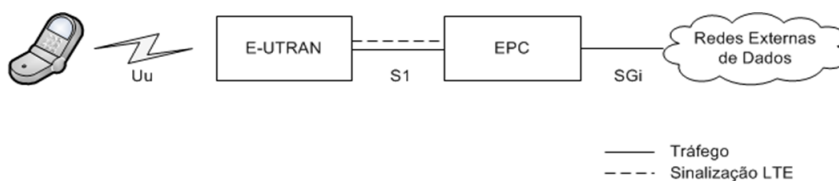


Figura 2.9: Visão de alto nível da arquitetura LTE

Dispositivos Móveis

O equipamento do usuário (*User Equipment*, UE), como em tecnologias anteriores, adota a tecnologia *SIM Card* (*Subscriber Identity Module Card*), um cartão inteligente que roda a aplicação *USIM* (*Universal Subscriber Identity Module*). Ele armazena dados específicos do usuário, tal como seu número de telefone e sua identidade na rede, e realiza diversos cálculos relacionados com a segurança, entre outras funcionalidades. O

LTE suporta celulares que estão usando USIM do *Release* 99 ou mais recente, mas não suporta o módulo de identidade do assinante (SIM), usado por versões anteriores do GSM [40] [14].

Foram especificadas no *Release* 8 cinco diferentes categorias de dispositivos. As principais exigências são o suporte a 64-QAM para *download* e diversidade de antenas. Para *upload*, 16-QAM é suportado para todas as categorias 1 a 4, sendo para a 5 a modulação 64-QAM suportada. Com exceção da categoria mais simples, todos os dispositivos devem suportar transmissão *multistream* de dados (MIMO), na qual vários canais de dados são transmitidos de forma paralela para o mesmo aparelho, sendo o número de canais determinado pelo total de antenas [39].

LTE suporta celulares que estão usando IP versão 4 (IPv4), o IP versão 6 (IPv6) ou a pilha dupla IP versão 4/6. O dispositivo recebe um endereço IP para cada rede de pacote de dados que está se comunicando [40] [14].

Visando a universalidade, os aparelhos dos usuários devem suportar também GSM (*Global System for Mobile Communications*), GPRS (*General Packet Radio Services*), EDGE (*Enhanced Data Rates For GSM Evolution*) e UMTS (*Universal Mobile Telecommunication System*). As interfaces e protocolos são projetados de tal maneira que uma conexão pode ser movida entre qualquer uma das tecnologias, caso o usuário saia da área de cobertura de LTE. A implantação da tecnologia começou independente das redes GSM e UMTS existentes, e sua integração com elas será feita assim que a tecnologia for suficientemente madura. LTE também suporta *roaming* com CDMA (*Code Division Multiple Access*), ou seja, o usuário pode ser transferido entre os dois tipos de rede mantendo o mesmo endereço IP e a estabilidade da sua conexão [39].

E-UTRAN, *Evolved UMTS Terrestrial Radio Access Network*

A base de toda a rede é o eNodeB, nome dado às estações base LTE. Suas principais funções são o envio de transmissões de rádio aos terminais móveis (*download*) e a recepção de dados dos mesmos (*upload*), dispondo de funções de processamento analógico e digital da interface LTE. A eNodeB também controla as operações de baixo nível dos seus dispositivos móveis [40] [14] [39].

Agindo como unidades autônomas, as estações eNodeB assumem o papel da central de controle RNC (*Radio Network Controller*) da tecnologia 3G UMTS, reduzindo a latência da rede. Assumindo o papel dessas centrais, essas unidades também são responsáveis por [40] [43]:

- Gerência de recursos de rádio;
- Gerência de usuários;
- Programação e transmissão de mensagens de *paging* e informações de *broadcast*;

- Garantia de QoS de acordo com o perfil do usuário;
- Gerência de mobilidade (*handovers*);
- Gerência de interferência.

A interface entre o eNodeB e o núcleo da rede é chamada S1, implementada via *link* de alta velocidade, pois requer taxas entre Mbits/s a Gbits/s para o transporte de dados de todos os usuários mais dados internos da rede [40] [39].

As estações base geograficamente próximas podem estar conectadas diretamente via interface X2, como mostra a Figura 2.10. Através dela as estações base trocam informações em operações de *handover* e coordenam ações na gerência de interferência. Caso não haja essa ligação, a conexão com o núcleo da rede é utilizada para realizar a troca de dados necessária. As relações entre estações base vizinhas podem ser previamente configuradas, ou as próprias estações podem detectar suas vizinhas, através de informação enviada pelos dispositivos móveis, o chamado ANR (*Automatic Neighbor Relation*) [40] [14] [39].

A comunicação pode ser logicamente dividida em planos de usuário e de controle. A primeira é responsável pelo transporte dos dados de usuário, encapsulados num *link* IP via protocolo GTP (*General Packet Radio Service Tunneling Protocol*), da mesma maneira do GPRS. Nesse processo, somente o endereço IP na camada de rede é alterado, enquanto que o IP do usuário permanece o mesmo. Os protocolos das camadas física e de enlace não são descritos na especificação LTE, possibilitando o uso de qualquer protocolo para o transporte de pacotes IP [39].

O protocolo do plano de controle é utilizado para transporte de dados de controle da rede. É por ele que a estação base comunica com o núcleo da rede, recebe dados de configuração, mantém a conexão do usuário e envia mensagens de sinalização sobre os usuários do sistema. Na pilha do protocolo tem-se o IP na base, logo acima o SCTP (*Stream Control Transmission Protocol*), protocolo equiparável ao TCP (*Transmission Control Protocol*) ou ao UDP (*User Datagram Protocol*), utilizados em tecnologias anteriores. O SCTP garante um grande número de conexões de sinalização independentes e simultâneas transportadas em sequência e faz a gestão de congestionamento e controle de fluxo [39].

As conexões lógicas S1 e X2 são independentes das tecnologias de transporte empregadas e são implementadas por uma rede interna IP, pela qual os dados são roteados e entregues, como mostrado na Figura 2.10. Cada elemento da E-UTRAN e da EPC possui um endereço IP, que permite a troca de dados e mensagens de sinalização [40] [14].

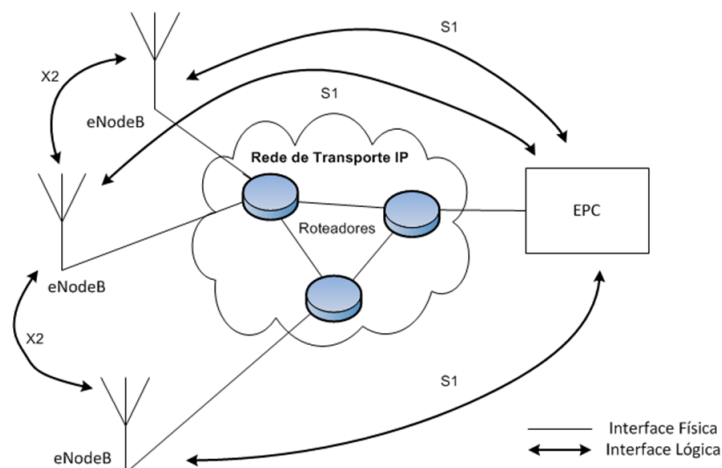


Figura 2.10: Arquitetura Interna da Rede de Transporte E-UTRAN [40]

EPC, Evolved Packet Core

O EPC é responsável pela comunicação com redes externas de dados, como Internet, redes corporativas ou sistemas de multimídia IP. Seus principais componentes são o *Home Subscriber Server* (HSS), *Mobility Management Entity* (MME), *Serving Gateway* (S-GW), e *Gateway PDN* (*Packet Data Network*). O HSS é a base de dados central que contém as informações de todos os assinantes da operadora de rede, e é um dos poucos componentes herdados do GSM e UMTS para LTE [40] [14].

O *Mobility Management Entity* (MME) é o elemento chave de controle da rede LTE. Gerencia os outros elementos da EPC através de mensagens de sinalização e as operações de alto nível dos dispositivos móveis da sua área de cobertura. Segundo [39], dentre suas várias funções, pode-se destacar:

- Autenticação de usuários e estabelecimento de barreiras;
- Gerenciamento de mobilidade NAS (*Non-Access Stratum*);
- Suporte a *handovers*;
- Interoperabilidade com outras redes de rádio;
- Suporte a serviços não baseados em IP.

O *gateway PDN* (*Packet Data Network*) provê conectividade entre os elementos da rede LTE e qualquer rede de dados externa, agindo como nó de entrada e saída do tráfego de dados do usuário. É este *gateway* o responsável pelo endereçamento IPv4 e/ou IPv6 dos dispositivos móveis realizado após a autenticação. Um dispositivo móvel pode ter conectividade simultânea com mais de um PDN para acesso a múltiplas redes de dados [40] [43].

O endereçamento IPv4, por dispor de uma menor gama de endereços a serem distribuídos, dispõe do uso de NAT (*Network Address Translation*), onde os usuários

são indiretamente endereçados, obrigando que todo o fluxo de dados passe por este nó concentrador que obtém o IP externo a rede [39]. A vantagem dessa configuração é que, por ter nas mãos todo o fluxo de dados, o PDN pode descartar pacotes maliciosos, como vírus de sondagem de rede, evitando uma carga desnecessária sobre a rede e o usuário [43].

O *Serving Gateway* (S-GW) faz um roteamento de alto nível, encaminhando pacotes de dados entre o *gateway* PDN e as estações base de uma determinada região geográfica. Atua também como âncora de mobilidade durante *handovers* inter-eNodeBs e com outras tecnologias 3GPP [40] [43].

A conexão do usuário com a Internet é feita por dois túneis independentes criados pelo MME, o primeiro é formado na rede de rádio e é alterado a cada *handover* realizado pelo usuário. Já o segundo, no núcleo da rede, é alterado somente se o *handover* for realizado entre duas estações base de MME diferentes [14].

O S-GW e MME são definidos independentemente, apesar de atuarem juntos, possibilitando que suas funções sejam executadas em nós diferentes da rede. Isso permite uma evolução individual de ambos, atendendo de forma adequada as demandas da rede, seja por sinalização adicional ou aumento capacidade de roteamento [39].

2.3.3 Arquitetura da Interface de Rádio LTE

A Figura 2.11 mostra a pilha de protocolos da interface de rádio LTE, destacando o fluxo de informações trocadas entre as camadas. No plano de usuário, as aplicações que rodam nos dispositivos móveis criam pacotes de dados a serem processados por protocolos tais como TCP, UDP e IP [40] [14].

Já no plano de controle, o MME se comunica com os dispositivos móveis utilizando protocolos da camada *Non-Access Stratum* (NAS). Nela, o MME dispõe dos recursos de gerência de mobilidade, pelo protocolo EMM (*Evolved Packet System Mobility Management*); gestão de sessões, pelo protocolo ESM (*Evolved Packet System Session Management*); além de realizar procedimentos de segurança. Essas mensagens são passadas para a camada RRC (*Radio Resource Control*), que escreve as mensagens de sinalização que serão trocadas com os terminais móveis [14].

Tanto no plano de controle quando no plano de usuário, as informações são entregues ao PDCP (*Packet Data Convergence Protocol*), seguindo um fluxo comum na seguinte ordem [40] [14]:

- *Packet Data Convergence Protocol* (PDCP);
- *Radio Link Control* (RLC);
- *Medium Access Control* (MAC);
- Camada física.

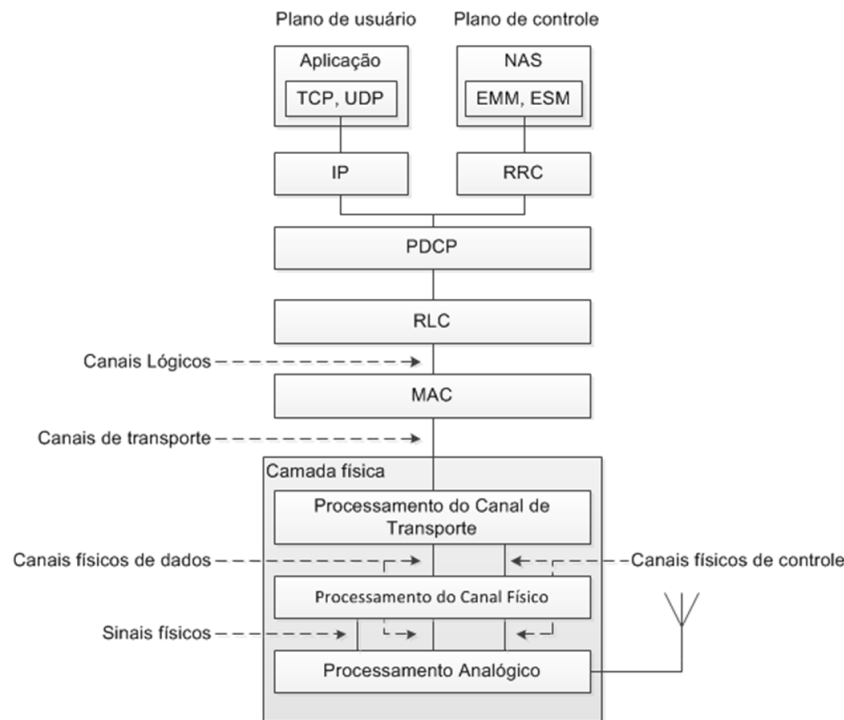


Figura 2.11: Pilha de protocolos da arquitetura de rádio LTE [40]

O *Packet Data Convergence Protocol* (PDCP) tem como suas principais funções a compressão e descompressão do cabeçalho IP, proteção, codificação e decodificação de pacotes. A compressão é baseada no algoritmo *Robust Header Compression* (ROHC), e visa a redução do número de bits a serem transmitidos, enquanto a proteção e codificação visam a manutenção da integridade dos dados [40] [15].

Radio Link Control (RLC) fica localizado na estação eNodeB, e oferece serviços para PDCP na forma de rádio portadores. Há sempre uma entidade RLC para um rádio portador configurado para o terminal móvel. É responsável pela segmentação e concatenação, manuseio de retransmissão, e a entrega dos pacotes de dados em sequência para as camadas mais altas [15].

De acordo com a decisão do escalonador, certa quantidade de pacotes IP, referidos como RLC SDU (*Service Data Unit*), é selecionada do *buffer* para transmissão. Eles são então segmentados ou concatenados para o tamanho adequado para transmissão, formando agora pacotes denominados RLC PDU (*Protocol Data Unit*). Tendo em vista as variadas possibilidades de taxas de transmissão que podem ser alcançadas, o tamanho dos PDUs varia dinamicamente, sendo grande para altas taxas e baixo para taxas menores [40] [15].

Em todo PDU é incluído um cabeçalho que, entre outras coisas, inclui um número sequencial, utilizado para entrega de pacotes na sequência correta. É através dessa numeração que o RLC exerce um papel importante na entrega de dados livre de erros. Sabendo a sequência de transmissão é possível identificar pacotes que faltam e requisitar

suas retransmissões [15].

Canais Lógicos e de Transporte

Os dados fluem entre as diferentes camadas de protocolo pelos seguintes canais [40]:

- Canais lógicos: entre RLC e MAC;
- Canais de transporte: entre MAC e a camada física;
- Canais físicos: dentro dos diferentes níveis da camada física.

Os canais lógicos podem ser distinguidos tanto pelo tipo de informação que carregam (plano de usuário ou de controle), como para quem entregam (canais dedicados ou compartilhados). Vários canais são definidos, como segue [40] [14] [15]:

- *Dedicated Traffic Channel* (DTCH): transmissão de dados de tráfego de um usuário específico;
- *Dedicated Control Channel* (DCCH): transmissão de informações individuais de controle entre um dispositivo móvel específico e a célula correspondente;
- *Broadcast Control Channel* (BCCH): transmissão de informação gerais da rede para todos os terminais móveis, informando aos dispositivos as configurações das células e os dados necessários para conexão;
- *Paging Control Channel* (PCCH): transmissão de mensagens *paging*, destinadas aos dispositivos em estado ocioso;
- *Common Control Channel* (CCCH): carrega mensagens úteis para usuários que estão estabelecendo conexão, em transição de status de ocioso para conectado;
- *Multicast Control Channel* (MCCH): transporte de informações de controle em transmissões *multicast*;
- *Multicast Traffic Channel* (MTCH): transporte de dados de tráfego em transmissões *multicast/broadcast*.

A camada *Medium Access Control* (MAC) faz a multiplexação de canais lógicos em canais de transporte, retransmissão *Hybrid ARQ* (*Automatic Repeat Request*) e escalonamento de *download* e *upload*. Da camada física, a camada MAC usa os serviços na forma de canais de transporte, definindo como e com quais características a informação é transmitida pela interface de rádio. Nesse nível, os dados são organizados em blocos, e cada bloco tem a duração de um intervalo de transmissão OFDM [15].

Os seguintes canais de transporte são definidos para LTE [40] [15]:

- *Downlink Shared Channel* (DL-SCH) e *Uplink Shared Channel* (UL-SCH): carregam a grande maioria dos dados, e são utilizados respectivamente para *download* e *upload*;

- *Broadcast Channel* (BCH): transmissão da maioria das informações do BCCH;
- *Paging Channel* (PCH): transmissão de informações do PCCH. Suporta recepção descontínua (DRX, *Discontinuous Reception*), isto é, o dispositivo móvel só fica ativo em instantes predefinidos para recepção, economizando bateria;
- *Multicast Channel* (MCH): suporte a transmissão *multicast/broadcast*.

A multiplexação dos canais lógicos em canais de transporte apropriados, e destes em canais físicos, pode ser visualizada na Figura 2.12 para *download* e na Figura 2.13 para *upload*.

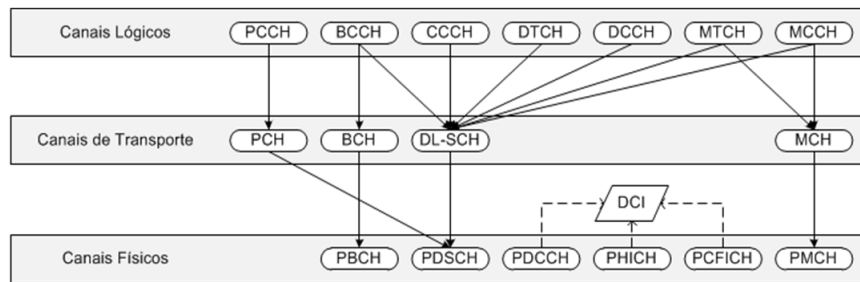


Figura 2.12: Mapeamento de Canais de Download [40]

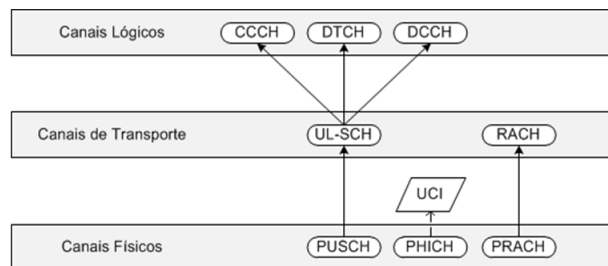


Figura 2.13: Mapeamento de Canais de Upload [40]

As principais diferenças entre os canais de transporte se encontram em suas abordagens de gestão de erro. Os canais de DL-SCH e UL-SCH são os únicos canais de transporte que usam as técnicas de pedido de retransmissão, além de serem capazes de adaptar a sua taxa de codificação para as alterações da SINR. Os demais têm taxa de codificação fixa, e fazem uso da *forward error correction* como técnica de tratamento de erro [40] [14].

Para a transmissão de mensagens de *paging* pelo PCH, a adaptação dinâmica dos parâmetros de transmissão pode, em certa medida, ser utilizada. Em geral, o processamento é semelhante ao tratamento de DL-SCH genérico. O MCH é usado para transmissões *multicast*, tipicamente com a operação da rede de frequência única, por transmitir a partir de várias células nos mesmos recursos com o mesmo formato e simultaneamente. Assim, a programação de transmissões MCH deve ser coordenada entre as células envolvidas e seleção dinâmica de parâmetros de transmissão pelo MAC não é possível [15].

Os canais de transporte de *downlink* restantes são baseados no mesmo processamento de camada física geral como o DL-SCH, embora com algumas restrições no conjunto de recursos utilizados. Para a transmissão de informações do sistema no BCH, um terminal móvel deve ser capaz de receber este canal de informação como um dos primeiros passos antes de acessar o sistema. Consequentemente, o formato da transmissão deve ser conhecido para os terminais, e não existe nenhum controle dinâmico de qualquer dos parâmetros de transmissão a partir da camada MAC neste caso [40] [15].

O processador de canal de transporte gera vários tipos de informações de controle para apoiar a operação de baixo nível da camada física, dentre elas podemos destacar dois tipos: *Uplink Control Information* (UCI) e *Downlink Control Information* (DCI). UCI contém vários campos, carregando informações para processamento ARQ, do SINR em função da frequência (*Channel Quality Indicator*, CQI), além de informações para multiplexação espacial. DCI contém a maior parte dos campos de controle de *download*, fazendo a manutenção da conexão terminal-estação base [40] [14].

Canal Físico

A camada física é dividida em três partes: o processado do canal de transporte, o processador de canal físico e o processador analógico.

O processador do canal de transporte se encarrega da gerência de procedimentos de tratamento de erro. Este também provê suporte às operações de manutenção de baixo nível da camada física, através do envio de informações de controle na forma de canais físicos de controle e sinais físicos, dados esses invisíveis para as camadas superiores [40] [14].

O processador de canal físico é a parte intermediária, e é ela que aplica as técnicas de modulação OFDMA (*Orthogonal Frequency-Division Multiple Access*) e SC-FDMA (*Single-Carrier Frequency-Division Multiple Access*) e, se for o caso, faz a transmissão via múltiplas antenas (MIMO).

O último, o processador analógico, converte a informação para forma analógica, filtra e envia pelo *link* de rádio para que esta seja transmitida [40] [14].

O canal físico corresponde a um conjunto de blocos de recursos dentro da grade tempo versus frequência, utilizado para a transmissão de um canal de transporte particular. Eles distinguem-se pelas formas em que o processador de canal físico os manipula, e pelas formas pelas quais são mapeados para os símbolos e as subportadoras utilizadas por OFDMA [40] [14] [15].

Há canais físicos sem um canal de transporte correspondente, conhecidos como canais de controle L1/L2. São utilizados para informações de controle de *download* e *upload*, provendo informações adicionais ao terminal e à estação base para operações envolvidas na transmissão de dados [15].

Os canais físicos definidos para LTE são os que seguem [40] [15]:

- *Physical Downlink Shared Channel* (PDSCH): principal canal, responsável por carregar dados e sinalização do DL-SCH e de informações de *paging* do PCH;
- *Physical Uplink Shared Channel* (PUSCH): transmissão de dados e sinalização do UL-SCH e de informações de controle de *upload* (*Uplink Control Information*, UCI);
- *Physical Broadcast Channel* (PBCH): carrega parte da informação do sistema, requerido para acesso de terminais à rede;
- *Physical Multicast Channel* (PMCH): utilizado em transmissões *broadcast*;
- *Physical Downlink Control Channel* (PDCCH) e *Physical Uplink Control Channel* (PUCCH): transmissão de informações de controle de *download* e *upload*, respectivamente;
- *Physical Hybrid-ARQ Indicator Channel* (PHICH): carrega informação de retransmissão, indicando para o terminal a retransmissão ou não do bloco de transporte;
- *Physical Control Format Indicator Channel* (PCFICH): fornece aos terminais as informações necessárias para decodificar o conjunto de PDCCHs;
- *Physical Random Access Channel* (PRACH): utilizado para acesso randômico.

Para transportar blocos no DL-SCH é adicionado um CRC (*Cyclic Redundancy Check*), usado para detecção de erro no receptor, seguido por *Turbo Coding*, para correção de erro. É definida uma taxa de correspondência para combinar o número de bits codificados com os recursos alocados para transmissão do DL-SCH, e os bits são então modulados para QPSK, 16-QAM ou 64-QAM. Segue-se o processo de mapeamento de antena, configurado para proporcionar diferentes regimes de transmissão multi-antena, incluindo diversidade de transmissão, *beam-forming* e multiplexação espacial. Finalmente, a saída é mapeada para os recursos físicos utilizados para o DL-SCH. Os recursos, assim como o tamanho de transporte de blocos e do esquema de modulação, estão sob o controle do escalonador [40] [15].

O procedimento da camada física para UL-SCH é o mesmo do DL-SCH, diferenciando do fato de que, por não suportar multiplexação espacial, não há mapeamento de antena no *upload* [15].

Na camada física tem-se os canais físicos de controle, cada qual carregando um tipo de informação de controle na forma de sinais físicos. Há vários tipos de sinais, que levam informações utilizadas para estimação de canal, escalonamento e sincronização, entre outras operações [11].

2.3.4 Esquemas de Transmissão LTE

Download e Upload

O LTE utiliza um esquema assimétrico de múltiplo acesso no *download* e *upload*, baseados em OFDMA (*Orthogonal Frequency-Division Multiple Access*) para *download* e SC-FDMA (*Single-Carrier Frequency-Division Multiple Access*) para *upload*.

O *download* é baseado em OFDMA com inserção de prefixo cíclico e com subportadoras de 15 kHz. O tamanho do prefixo cíclico é escolhido de tal maneira que suporte o máximo atraso de propagação do canal de rádio [40] [43].

A transmissão no *upload* é feita via SC-FDMA, também com inserção de prefixo cíclico a fim de atingir ortogonalidade entre diferentes usuários e melhorar a eficiência na equalização nos receptores (estações base eNodeB). O alto grau de semelhança deste esquema com o adotado no *download* leva a uma uniformização e simplificação na implementação do esquema como o todo, pois permite a adoção dos mesmos parâmetros de transmissão. A modulação SC-FDMA foi adotada principalmente pela sua relação de potência de pico-a-média (PAPR), relativamente inferior a de OFDMA. Essa propriedade reflete diretamente no consumo de potência do transmissor. Como no caso de *uplink* o transmissor é um dispositivo móvel, o consumo de energia é um fator importantíssimo e deve ser reduzido ao máximo, levando em conta que o dispositivo tem energia disponível limitada [40] [43].

A Figura 2.14 apresenta a diferença básica entre OFDMA e SC-FDMA. Para um conjunto de símbolos a ser enviado, OFDMA o faz transmitindo-os em paralelo, cada um acompanhado por um prefixo cíclico. SC-FDMA transmite os símbolos sequencialmente, e insere o prefixo cíclico no final da transmissão. Sabendo que o período dos símbolos de ambos os esquemas é o mesmo, SC-FDMA requer N vezes mais largura de banda para obter a mesma taxa de transmissão de N símbolos OFDMA transmitidos paralelamente [43] [39].

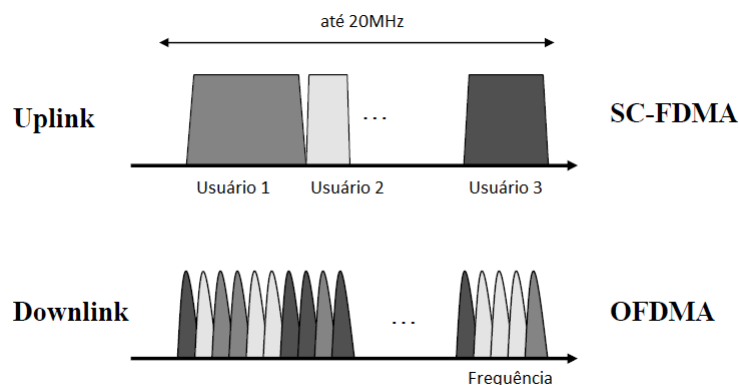


Figura 2.14: Comparação entre esquemas de transmissão LTE [25]

SC-FDMA exige processamento adicional ao transmitir o sinal, pois ao invés de dividir os bits em subportadoras diferentes, o sinal temporal é convertido para frequência distribuindo a informação de cada bit sobre todas as subportadoras que serão utilizadas na transmissão, reduzindo as diferenças de potência entre elas. O número de subportadoras utilizadas depende das condições do canal, da potência disponível e do número de usuários simultâneos no *upload* [43] [39].

Parâmetros de Transmissão LTE

As transmissões LTE são baseadas na unidade de tempo básica T_s , intervalo de amostragem num sistema de Transformada Rápida de Fourier com 2048 pontos. Seu valor é dado pela seguinte equação:

$$T_s = \frac{1}{2048 \cdot SCBW} \approx 32.6ns. \quad (2-3)$$

Nesta equação, $SCBW$ é o valor da largura de banda de uma única portadora, que no caso da tecnologia LTE tem valor 15000 Hz. A unidade de transmissão é um símbolo, também chamada elemento de recurso, e sua duração é dada por $2048 \cdot T_s = 66.667\mu s$. Um símbolo pode transmitir vários bits, dependendo do esquema de modulação e codificação adotados. Quanto melhor a condição de canal, mais bits de informação podem ser transmitidos no mesmo elemento de recurso. Condições excelentes permitem o uso de 64-QAM (6 bits/símbolo). Em condições de menor qualidade serão aplicadas modulações de menor grau, 16-QAM ou QPSK (4 bits/símbolo e 2 bits/símbolo, respectivamente) [40] [14] [39].

Os símbolos são agrupados em *slots* de 0.5 ms ($15360 \cdot T_s$). Um *slot* pode ser organizado internamente de duas maneiras. Em condições normais de transmissão, são carregados 7 símbolos e cada um é precedido de um prefixo cíclico de tamanho normal, com duração de $4.7 \mu s$ ($144 \cdot T_s$), exceto o primeiro símbolo, que tem prefixo de $5.2 \mu s$ ($160 \cdot T_s$) de duração, mais longo afim de ocupar o espaço restante da divisão. Esta configuração suprime interferência intersimbólica com atrasos de propagação de até $4.7 \mu s$, correspondente a células com raio de até 1.4 km [40] [14] [39].

Em células grandes ou com grande quantidade de interferência intersimbólica, o prefixo cíclico de tamanho normal pode não ser suficiente. Nesses casos, dispõe-se do prefixo cíclico estendido, com duração de $16.7 \mu s$ ($512 \cdot T_s$), utilizado em células de raio de até 5 km. Nessa situação cada *slot* conterà somente 6 símbolos [14].

Num nível mais alto, as transmissões de *download* e *upload* são organizadas em *subframes* (1 ms), que são pares de *slots*, e *frames*, que são agrupamentos de 10 *subframes* (10 ms). Quando uma estação base transmite para um terminal móvel, ela escalona transmissões via PDSCH num *subframe* por vez, mapeando cada bloco num

conjunto de subportadoras dentro do *subframe*. Procedimento similar acontece no *upload* [40] [14].

LTE suporta dois tipos de estrutura para *frames*: o Tipo 1, aplicável a *Frequency Division Duplex* (FDD), e o Tipo 2, aplicável a *Time Division Duplex* (TDD). No esquema FDD, *download* e *upload* são separados por frequência. Logo, dada uma subportadora, todos os seus *subframes* estão disponíveis para *upload* ou *download*. Os *frames* são numerados repetidamente de 0 a 1023, e ajudam no gerenciamento de processos de mudança mais lenta, tais como transmissão de informações do sistema e sinais de referência [40] [14] [43].

Para esquemas TDD, a separação entre sinais de *download* e *upload* é feita no tempo. Cada *subframe* pode ser alocado tanto para um como para outro tipo de transmissão, de acordo com as diferentes configurações que podem ser adotadas pelas células. Para diferenciar o tipo de informação carregada, lança-se mão de um *subframe* especial, de 1 *ms* de duração, contendo o final da transmissão de *upload*, um período de guarda e início dos *slots* de *download* [43].

A informação é organizada em uma grade de recursos que reflete a grade de transmissão tempo versus frequência OFDM. A Figura 2.15 mostra a grade de recursos para prefixo cíclico de tamanho normal. Os elementos de recurso são agrupados a cada 0.5 *ms* (um *slot*) por 180 kHz (12 subportadoras contíguas) [40] [43].

No domínio da frequência os recursos são agrupados em 12 subportadoras de 15KHz, totalizando um largura de banda de 180KHz. Um bloco de recurso (RB, *Resource Block*) é definido como uma unidade de 12 subportadoras durante um *slot* de tempo [1]. No sistema LTE os blocos de recursos são escalonados sempre em pares de RBs, chamados assim de blocos de escalonamento (SB), com duração de 1ms.

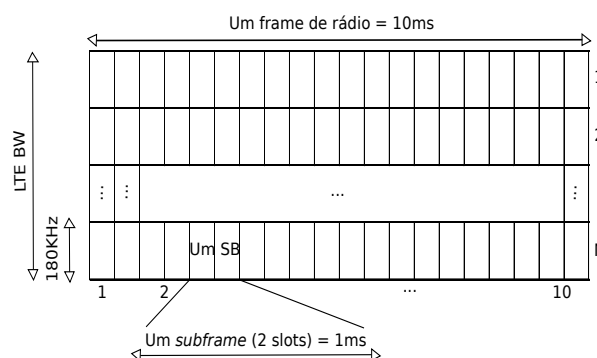


Figura 2.15: Estrutura básica de um frame LTE no domínio do tempo e frequência [45]

A transmissão pode ocorrer de forma localizada ou distribuída. No primeiro modo, denominado *Localized Virtual Resource Blocks* (LVRBs), a transmissão é feita por grupos formados por elementos vizinhos na grade OFDM tempo versus frequência. Esse tipo de transmissão requer um retorno do dispositivo móvel sobre a qualidade das

subportadoras utilizadas, permitindo à estação base selecionar o esquema de transmissão mais adequado. A segunda opção de transmissão é chamada *Distributed Virtual Resource Blocks* (DVRBs), onde os símbolos componentes de um bloco estão espalhados por todo o espectro. Nesse caso ou não há resposta do dispositivo móvel sobre a qualidade do canal, ou esta resposta faz referência ao canal de modo geral [40] [39].

A célula pode ser configurada para diferentes larguras de banda, opções estas listadas na Tabela 2.2. Pode-se notar que há sempre uma diferença entre a banda disponível e a ocupada. Essa diferença é preenchida pelas bandas de guarda nos limites superior e inferior da banda de frequência, minimizando a quantidade de interferência com a próxima banda. As duas bandas de guarda possuem geralmente a mesma largura, mas o operador de rede pode ajustá-la, se necessário, mudando a frequência central em unidades de 100 kHz [40] [14].

Banda Total (MHz)	Blocos de Recurso	Subportadoras	Banda Ocupada (MHz)	Banda de Guarda (MHz)
1.4	6	72	1.08	2×0.16
3	15	180	2.7	2×0.15
5	25	300	4.5	2×0.25
10	50	600	9	2×0.5
15	75	900	13.5	2×0.75
20	100	1200	18	2×1.00

Tabela 2.2: *Larguras de banda suportadas no LTE [14]*

Símbolos de Referência e de Sincronização

Nem todos os símbolos de um bloco de recurso são utilizados para transmitir dados. Para que os dispositivos móveis consigam detectar as subportadoras, diferenciar estações base e estimar a qualidade do canal, os chamados símbolos de referência, também conhecidos como sinais de referência, são incorporados em um padrão predefinido sobre a largura do canal inteiro. São definidos vários tipos, diferenciando-se pela sua funcionalidade e pelo padrão em que são transmitidos. São inseridos em cada sétimo símbolo sobre o eixo do tempo e em cada sexta subportadora no eixo de frequência [40] [39].

Para sincronização inicial são utilizados dois tipos de sinais adicionais, referidos como os sinais de sincronismo primário e secundário, transmitidos em todos os primeiros e sextos *subframes* a cada 72 subportadoras do canal. Em cada uma dessas subportadoras, um símbolo é utilizado para cada sinal de sincronização [39].

Há também os sinais de referência transmitidos no *upload*, classificados em dois tipos: de demodulação e de referência de sondagem. Os símbolos de demodulação são utilizados para a estimativa do canal no receptor, habilitando a eNodeB a demodular os canais de controle e de dados. São transmitidos no quarto símbolo de cada *slot*, e

compreendem a mesma largura de banda dos dados de *uplink* alocados. O outro tipo é o sinal de referência de sondagem, que fornece informações de qualidade do canal. Serve como base para as decisões de planejamento na estação base. São transmitidos no último símbolo do *subframe*, em diferentes partes do espectro onde não há dados de *upload* disponíveis [40] [43].

Alocação de Recursos em Redes LTE

Este capítulo apresenta conceitos sobre alocação de recursos possíveis em redes LTE, onde estes são alocados dinamicamente aos diferentes usuários com base na qualidade do canal de transmissão e em diferentes parâmetros de QoS. É apresentada a modelagem multifractal β MWM, utilizada para estimativa da banda efetiva do tráfego de rede. Também são apresentados dois métodos de alocação dos recursos, o algoritmo com garantia de QoS [24] e o algoritmo baseado em *Particle Swarm Optimization* (PSO) [42]. Estes dois algoritmos são utilizados como referência para o estudo comparativo com os algoritmos propostos nos capítulos posteriores.

3.1 Alocação de Recursos no *Download* e *Upload*

As redes LTE tem como grande desafio a comunicação multiusuário, vários terminais móveis na mesma região a serem atendidos com altas demandas de taxas de transmissão e baixo *delay*, em uma largura de banda finita. Para o atendimento dessas expectativas, adota-se a técnica de múltiplo acesso OFDMA, na qual se compartilha uma largura de banda com os usuários, destinando a cada um uma fração dos recursos totais disponíveis e acomodando diferentes requisitos de taxas de dados e QoS [40] [47].

A transmissão de dados em LTE consiste então no compartilhamento de um canal físico constituído de uma grade tempo versus frequência de blocos de recursos (*slots*), unidades mínimas de alocação, que são dinamicamente distribuídos entre os usuários. Há uma forte motivação na otimização do escalonamento e distribuição de recursos, pois implicam diretamente na melhora do desempenho do sistema através do aumento da eficiência espectral da interface sem fio [40] [47].

O escalonador é parte da camada MAC e controla a distribuição dessas unidades de recursos no *download* e *upload*. No *download* controla-se dinamicamente para quais terminais transmitir e qual a quantidade de recurso destinada a cada um deles. A cada intervalo de *subframe*, o escalonador aloca blocos de recursos para transmissão pelo canal DL-SCH para o dispositivo móvel, informação esta utilizada para processamento na camada física [40] [15].

O escalonamento para *upload* é feito de forma independente do *download*. É nessa fase que se determinam também para quais terminais móveis transmitir nos canais UL-SCH e a distribuição da grade de recursos entre usuários. A estação base determina a carga útil a ser transmitida pelo terminal e o formato de transporte a ser aplicado, tomando por base o terminal e não a portadora utilizada. O dispositivo móvel tem autonomia no sentido de selecionar a partir de qual portadora de rádio os dados são obtidos, manipulando autonomamente a multiplexação de canais lógicos [40] [15].

Para grandes larguras de banda com quantidade significativa de seletividade em frequência, e especialmente em usuários com baixa velocidade, a possibilidade de adaptar a transmissão à condição do canal traz enormes benefícios. Essa variabilidade aleatória da condição do canal proíbe o uso contínuo de uma única modulação, adotando-se assim a estratégia de codificação e modulação adaptativa (*Adaptive Modulation and Coding*, AMC) [40] [47].

Os fatores que limitam a taxa de dados são a largura de banda e a relação sinal ruído. A estratégia AMC traz uma ideia simples e eficiente, qual seja, transmitir a taxas maiores quando o canal tem alto SINR, e transmitir a menores taxas para menores valores de SINR, no intuito de evitar perda excessiva de pacotes. Baixas taxas são alcançadas pelo uso de constelações pequenas, como por exemplo QPSK, acompanhado de técnicas de correção mais robustas. E como esperado, taxas maiores são conseguidas com constelações maiores tais como 16 e 64QAM, e procedimentos de correção de erros menos robustos [47] [40] [15].

O escalonamento chamado de *Channel-dependent scheduling* é principalmente utilizado para download [15]. A fim de dar suporte a esse procedimento, o terminal móvel reporta periodicamente um status do canal para a estação base, informando a qualidade instantânea do canal no domínio do tempo e da frequência [40] [15].

O *Channel Quality Indicator* (CQI) é um indicador de quatro bits, e tem como função apontar a taxa de transmissão máxima que um dispositivo móvel pode alcançar com uma taxa de erro por bloco de até 10%. Ele depende principalmente do SINR, mas também varia com a capacidade do *hardware* do receptor. A Tabela 3.1 mostra como a estação base interpreta o CQI, no intuito de aplicar o esquema de modulação e taxa de codificação apropriados [40] [14].

A modulação de baixa ordem, por exemplo, QPSK, é mais robusta e pode tolerar altos níveis de interferência, mas fornece uma menor taxa de transmissão de bits. A modulação de alta ordem, por exemplo, 64-QAM, oferece uma taxa de bits maior, mas é mais propensa a erros por conta de sua maior sensibilidade a interferência, ruído e erro de estimação de canal. Por isso a modulação de alta ordem é útil apenas quando a SINR é suficientemente elevada [14]. Para uma dada modulação, a taxa de codificação pode ser escolhida dependendo das condições do enlace de rádio. Uma taxa de codificação inferior

pode ser utilizada em condições ruins de canal e uma taxa de codificação mais elevada, no caso de SINR alta [14].

CQI	Modulação	Taxa de Codificação (1/1024)	Bits de informação por símbolo
0	-	0	0.00
1	QPSK	78	0.15
2	QPSK	120	0.23
3	QPSK	193	0.38
4	QPSK	308	0.60
5	QPSK	440	0.88
6	QPSK	602	1.18
7	16-QAM	378	1.48
8	16-QAM	490	1.91
9	16-QAM	616	2.41
10	64-QAM	466	2.73
11	64-QAM	567	3.32
12	64-QAM	666	3.90
13	64-QAM	772	4.52
14	64-QAM	873	5.12
15	64-QAM	948	5.55

Tabela 3.1: CQI e respectivos esquemas de modulação e taxa de codificação [14]

A possibilidade de uma distribuição de recursos dependente da condição do canal, observando-se tanto as variações temporais quanto na frequência, traz enormes benefícios. Um exemplo do funcionamento desse modelo pode ser visualizado na Figura 3.1. Para serviços sensíveis a atrasos, em decisões tomadas atentando-se somente a variações no domínio do tempo, os recursos podem ser forçados a um usuário em particular, apesar de a qualidade do canal não estar no seu auge. Em tais situações, explorando as variações de qualidade de canal também no domínio da frequência, pode-se melhorar o desempenho global do sistema. O esquema de transmissão adotado na tecnologia LTE explora bem essas variações, pois além de observar tempo e frequência, o faz analisando bloco de baixa granularidade, a cada 1 *ms* no tempo e 180 kHz na frequência [40] [15].

Embora a estratégia de escalonamento e distribuição de recursos seja específica da implementação e não padronizada pela 3GPP, o principal objetivo é sempre tomar vantagem da variabilidade de condição de canal em tempo e frequência entre os dispositivos móveis, na tentativa de atender três principais requisitos do sistema: eficiência espectral, *fairness* (índice de justiça) e QoS [40] [47].

É uma tarefa difícil atender a esses requisitos simultaneamente e de maneira otimizada. Os esquemas que priorizam maximização da taxa de dados são injustos, pois

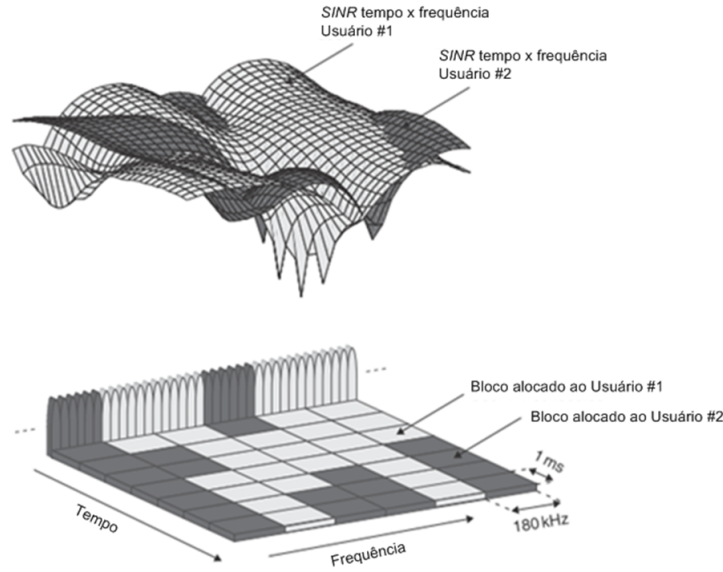


Figura 3.1: *Channel-dependent scheduling no download para dois usuários, nos domínios do tempo e da frequência [40]*

priorizam dispositivos móveis em melhores condições, já um nível maior de justiça pode levar a uma ineficiência no uso da banda disponível [47].

Soluções ótimas podem ser alcançadas alocando-se sempre para um determinado usuário os subcanais disponíveis que possuem melhores condições. O procedimento torna-se complicado ao se levar em conta situações nas quais o melhor subcanal de um usuário pode ser também o melhor de outro, e que este não tenha outros subcanais de boa qualidade. Diferente das visões tradicionais onde o desvanecimento é considerado somente um prejuízo, aqui ele atua também como um seletor aleatório de canal, aumentando a diversidade multiusuário [40] [47].

3.1.1 Estratégias de Alocação de Recursos no *Download*

Uma estação eNodeB serve múltiplos usuários através de múltiplos canais compartilhados, regulando a transmissão de dados em ambas as direções. Todos os pacotes das camadas mais altas destinados a terminais móveis são classificados na eNodeB em fluxos de dados com diferentes requisitos de QoS. Cada fluxo é associado com um portador, que pode ser considerado como uma conexão entre o *gateway* PDN e o usuário móvel. Assim que classificado pela camada MAC, os fluxos são designados a *slots* na grade de transmissão tempo versus frequência, isto é, segmentam-se os dados depois da modulação em blocos que preenchem um *slot* OFDMA [40] [47].

Cada usuário tem uma fila de pacotes de dados a receber. Cabe então à estação base servir as demandas simultaneamente de vários usuários, alocando um ou mais *slots* dos disponíveis na grade de recursos. Nesse modelo, uma grande variedade de taxas de

dados pode ser suportada, explorando-se o conhecimento das condições do canal sem fio e as demandas de cada usuário [40] [47].

Segundo [47], das várias razões para o estudo e otimização da distribuição de recursos no cenário LTE, pode-se destacar os benefícios do uso da diversidade de usuário na maximização da saída do sistema, tarefa que deixa de ser simples pela necessidade de atendimento dos parâmetros de QoS do sistema, em especial o critério de justiça. A especificação LTE não determinou a forma como deve ser feita essa alocação, deixando livre o desenvolvimento de procedimentos inovadores que determinem como distribuir a banda disponibilizada.

Pode ser observado que as estratégias de alocação de recursos de OFDM são principalmente limitadas pelo nível das subportadoras, e todas subportadoras alocadas a um usuário específico devem adotar o esquema de modulação e codificação que maximize a eficiência espectral, isto é, utilize a maior taxa de informação por símbolo para seu nível de SINR. Dentro de sistemas LTE, entretanto, os recursos são alocados em unidades de recursos denominados *Scheduling Block* (SB), e todos os SBs designados a um determinado usuário devem utilizar o mesmo esquema de modulação e codificação (*Modulation and Coding Scheme*, MCS) num dado intervalo de transmissão (*Transmission Time Interval*, TTI) [40] [15].

3.2 Modelo do Sistema e Formulação do Problema

Considerando um cenário de transmissão *downlink* de um sistema LTE de uma antena, com N blocos de recursos disponíveis por TTI (*Transmission Time Interval*), quantidade de potência distribuída igualmente entre todas as subportadoras e K usuários servidos a taxas mínimas R_k . A taxa mínima R_k é definida por valores fixos ou valores estimados adaptativamente, conforme será explicado na Seção 3.5. Um bloco de recurso é definido como N_s símbolos OFDM consecutivos no domínio da frequência e N_{sc} subportadoras no domínio do tempo. Considerando que existem sinais pilotos e de controle nos blocos de recursos, apenas $N_{sc}^d(s)$ subportadoras podem ser utilizadas para transferência de dados no s -ésimo símbolo OFDM, onde $s \in \{1, 2, \dots, N_s\}$ e $N_{sc}^d(s) \leq N_{sc}$. Seja $R_j^{(c)}$ a taxa de código associada ao MCS (*Modulation and Coding Scheme*) $j \in \{1, 2, \dots, J\}$, onde J é o número total de MCS suportados, M_j o tamanho da constelação do MCS j e T_s o tempo do símbolo OFDM, a taxa de bits de um bloco de recursos $r^{(j)}$ alcançada para o MCS j é dada por [45]:

$$r^{(j)} = \frac{R_j^{(c)} \log_2(M_j)}{T_s N_s} \sum_{s=1}^{N_s} N_{sc}^d(s). \quad (3-1)$$

O indicador de qualidade de canal (*Channel Quality Indicator*, CQI) é definido no LTE em termos da taxa de código e esquema de modulação, e tem a informação de qual MCS deve ser adotado para o usuário k no bloco de recurso (*Scheduling Block*, SB) n .

Cada SB é alocado a apenas um usuário em um TTI. Seja $q_n(i) \in \{1, 2, \dots, K\}$ o usuário alocado no bloco de recurso n no TTI i . A taxa de bits nesse bloco de recurso depende do MCS do usuário, logo a taxa de bits do mesmo bloco de recurso pode ser diferente para cada usuário. Alocar os blocos de recursos aos usuários com maiores taxas de bits tende a aumentar a utilização da rede.

Seja $x(t) = [x_1(t), x_2(t), \dots, x_N(t)]$ o vetor de tamanho N , composto pelos elementos $x_n(t)$, $n \in N$, que associa cada bloco de recurso a um usuário no instante de tempo t , por exemplo, se $x_1(t) = 2$, o bloco de recurso 1 é alocado para o usuário 2 no instante de tempo t ; $m_{k,n}(t)$ o MCS adotado para o usuário k no bloco de recurso n no instante de tempo t , a taxa de bits r_k do usuário k , no instante de tempo t é dada por [42] [45]:

$$r_k(t) = \sum_{n=1}^N I(x[n] = k) r^{(m_{k,n}(t))}, \quad (3-2)$$

onde $I(x[n] = k)$ é 1 se $x[n] = k$ e 0 caso contrário. A taxa de bits total T_b do sistema é [42] [45]:

$$T_b(t) = \sum_{k=1}^K r_k(t). \quad (3-3)$$

Maximizar a taxa de bits total do sistema T_b é uma forma de melhorar a utilização da rede. Porém, juntamente com o aumento da taxa de bits total do sistema é necessário atender certos requisitos de banda de cada usuário. Assim, tem-se um problema de otimização que consiste em maximizar a taxa do sistema atendendo a taxa mínima de cada usuário [42] [45]:

$$(x) : \max T_b, \quad (3-4)$$

sujeito a:

$$r_k \geq R_k \quad \forall k. \quad (3-5)$$

O problema de otimização é formulado então pela função-objetivo como sendo a equação 3-4, dado a taxa total alcançada por TTI, e as restrições são dadas pela equação 3-5, que garante o atendimento às taxas de dados mínimas dos usuários [24].

Analizando-se a equação 3-4, pode ser verificado que se trata de uma otimização não linear. Sua implementação por técnicas baseadas em gradiente pode não garantir uma otimização global e a complexidade crescente exponencialmente inviabiliza a aplicação

em tempo real. São propostos então métodos variados de alocação, tal como os métodos [24] e [42], descritos a seguir.

3.3 Algoritmo de Alocação com Garantia de QoS

O algoritmo de alocação com garantia de QoS surgiu da necessidade de reduzir a complexidade computacional do problema de otimização da vazão total descrito pela equação 3-4 atendendo a restrição da taxa mínima de transmissão dada pela equação 3-5 [24]. Os resultados apresentados em [24] mostram que o algoritmo QoS garantido apresenta valores de vazão satisfatórios e consegue atender ao critério de taxa mínima fixado no cenário simulado. Porém, as simulações que serão apresentadas nos capítulos 4 e 5 mostram que o algoritmo QoS garantido apresenta limitações em termos de taxa de perda e retardo.

O algoritmo é composto por dois passos [24]:

1. Estima o número de SBs necessários por usuário baseado na razão entre taxas mínimas exigidas por eles e suas respectivas médias dos ganhos de canal;
2. Em atendimento a restrição das taxas mínimas, aloca SBs para os usuários de acordo com suas prioridades.

As principais etapas do algoritmo são descritas a seguir [24].

3.3.1 Estimativa de SBs

Cálculo da média de ganho de canal de cada usuário

Sabendo que um único MSC é empregado em todos os SBs alocados para um determinado usuário, a condição de canal média é considerada no lugar da condição individual do canal para cada SB. Para reduzir a sobrecarga de retorno do status do canal, é adotado o método baseado no CQI *threshold* λ_k [24]. A cada alocação, o usuário retorna somente valores de CQI maiores que o limiar (*threshold*). Para aqueles valores que não são retornados para o usuário k , assume-se que seu CQI seja zero. Diferentes usuários podem ter diferentes valores de *threshold*. Por exemplo, para usuário no centro da célula, altos valores de *threshold* podem ser considerados, enquanto para usuário situados na borda, valores relativamente baixos são adotados. Sejam $\bar{g}_k$ e α_k , respectivamente, o ganho de canal médio do usuário k e o número de SBs com valores de CQI informados [24]. Pode-se dizer então que [24]:

$$\bar{g}_k = \frac{1}{\alpha_k} \sum_{n=1}^{\alpha_k} g_{k,n}, g_{k,n} \geq \lambda_k \text{ e } \alpha_k \leq N. \quad (3-6)$$

Estimação do número de SBs requeridos por cada usuário

Sabendo-se o ganho médio de canal de cada usuário, o número de SBs necessários é calculado pela razão entre a taxa mínima requisitada e $\bar{g}_k$. Sendo N_k o número de SBs alocados ao usuário k , N_k deve satisfazer as seguintes condições [24]:

$$\frac{R_1}{\bar{g}_1} : \frac{R_2}{\bar{g}_2} : \dots : \frac{R_K}{\bar{g}_K} = \phi_1 : \phi_2 : \dots : \phi_K, \quad (3-7)$$

$$\sum_{k=1}^K \phi_k = 1, \quad (3-8)$$

$$N_1 : N_2 : \dots : N_K = \phi_1 : \phi_2 : \dots : \phi_K. \quad (3-9)$$

O número de SBs obtidos por cada usuário k é proporcional a sua taxa mínima requerida e inversamente proporcional à média do ganho de canal. Essas condições são razoáveis desde que quanto melhor a condição do canal, mais dados podem ser transmitidos por unidade de recurso e menos recursos devem ser alocados para uma dada requisição de taxa [24].

Seja a função $\lfloor x \rfloor$ tal que retorne o maior número inteiro igual ou menor que x . Se $x = 0$, seu valor é assumido como 1. O número N_k de SBs alocados para cada usuário é obtido de:

$$N_k = \lfloor N_{\phi_k} \rfloor. \quad (3-10)$$

3.3.2 Alocação de SBs

Sabendo que um único MCS é adotado para todos SBs alocados para determinado usuário, quando estes SBs têm diferentes valores de CQI adota-se o MCS com menor CQI para assegurar que todos os dados transmitidos sejam corretamente recebidos. Isso resulta num decréscimo da taxa final que o usuário poderia alcançar [24].

O usuário com melhor condição média de canal e menor taxa mínima requisitada é designado primeiro, fazendo com que um maior MCS possa ser escolhido para o usuário e o resto dos SBs possam ser alocados aos demais, com alta probabilidade de alcançar seus requisitos de taxa [24].

A principal ideia do algoritmo de escalonamento com garantia de QoS pode ser expressa por [24]:

- Calcula as prioridades dos usuários e os ordena em ordem decendente. A prioridade p é definida como segue:

$$\text{se } \bar{g}_k > \bar{g}_i, \text{ então } \rho_k > \rho_i. \quad (3-11)$$

- Aloca os SBs para cada usuário.

De acordo com as prioridades já definidas, a alocação se dá usuário a usuário. Se o número de SBs estimados no primeiro passo for alocado e a taxa mínima exigida não for alcançada, mais SBs serão designados para este usuário até satisfazer sua requisição de taxa. Por outro lado, se os requisitos de todos os usuários forem atendidos e ainda houver SBs não alocados, estes serão designados para o usuário com maior prioridade [24].

A descrição detalhada do algoritmo de alocação de recursos de *download* para sistemas LTE com garantia de QoS é apresentada no Algoritmo 3.1.

Algoritmo 3.1: Algoritmo de escalonamento com garantia de QoS [24]

Entrada: K : número de usuários

N : número de SBs

R_k : taxa requisitada pelo usuário k

N_k : número de SBs estimados pelo usuário k

G : matriz $N \times K$ com o CQI dos usuários

Saída: r_k : vetor de taxas alcançadas por usuário k

1 **início**

2 $W = \{1, 2, \dots, N\}$

3 $S_k = \{\}, k \in \{1, 2, \dots, K\}$

4 $\rho_{k,n} = 0, k \in \{1, 2, \dots, K\}, n \in \{1, 2, \dots, N\}$

5 $r_k = 0, k \in \{1, 2, \dots, K\}$

6 Iteração de $k = 1$ até K

7 Se $W \neq \{\}$, vai para próximo passo. Caso contrário, vai para o final

8 Se $k < K$, aloca os SBs remanescentes para o usuário 1. Caso contrário, vai para próximo passo

9 Escolhe N_k SBs para o usuário k em concordância com o critério na condição que $W \neq \{\}$, e coloca os SBs nos grupos S_k de acordo

10 Encontra o SB com menor CQI no grupo S_k , então determina o maior MCS que pode ser utilizado pelo usuário k

11 Calcula r_k para o usuário k com o MCS escolhido no passo anterior

12 Se $r_k \geq R_k$, então $k = k + 1$ e vai para o passo 6. Caso contrário, vai para o próximo passo

13 Se $W \neq \{\}$, continua alocando um SB para o usuário k de acordo com o critério e coloca os SBs dentro dos grupos S_k de acordo. Vai para o passo 9

3.4 Algoritmo de Alocação Baseado em *Particle Swarm Optimization* (PSO)

O algoritmo de alocação baseado em otimização PSO (*Particle Swarm Optimization*) é uma das várias soluções propostas para solução do problema de otimização da vazão total descrito pela equação 3-4, tendo em vista a restrição da taxa mínima de transmissão dada pela equação 3-5 [42]. Os resultados apresentados em [42] mostram que o algoritmo PSO tem desempenho superior em termos de vazão total. Porém, como será apresentado nos capítulos 4 e 5, as simulações mostram que o algoritmo PSO apresenta limitações em relação ao critério de justiça, retardo e tempo de processamento.

A otimização PSO é uma heurística estocástica, sub-ótima, baseada em população e de fácil implementação. A população é chamada de enxame e cada indivíduo, que corresponde a uma solução para o problema, é chamado de partícula. Na otimização PSO padrão cada partícula possui posição, velocidade e memoriza a melhor posição da partícula encontrada até o momento, também chamada de melhor posição local. A melhor posição de partícula na população, ou seja, a solução com menor custo, também é memorizada. Os vetores velocidade e posição são variáveis contínuas [42].

Na inicialização cada partícula possui posição e velocidade aleatórias. O algoritmo procura a solução ótima através de atualizações das posições e velocidades de cada partícula, levando em conta as velocidades, as melhores posições das partículas e a melhor posição da população, até um critério de parada. As posições e as velocidades das partículas são atualizadas segundo as equações [42]:

$$v_{t+1} = wv_t + r_1c_1(P_t - X_t) + r_2c_2(G_t - X_t) , \quad (3-12)$$

$$X_{t+1} = X_t + v_{t+1} , \quad (3-13)$$

onde w é o peso de inércia, c_1 e c_2 são taxas de aprendizagem, r_1 e r_2 são dois números aleatórios gerados segundo uma distribuição uniforme $[0,1]$, v_t , X_t e P_t são, respectivamente, a velocidade, a posição, a melhor posição da partícula no instante de tempo t , e G_t é a melhor posição da população no instante de tempo t [42].

O Algoritmo 3.2 apresenta o funcionamento da otimização PSO.

Algoritmo 3.2: Algoritmo *Particle Swarm Optimization* [42]

início

Passo 1: Inicialização:

Inicializa as posições $X(i)$ e velocidades $v(i)$ de cada partícula i com valores aleatórios

Define melhor posição da partícula $P(i) = X(i)$, onde i é o índice da partícula

Calcula o custo $C(i)$ de cada partícula i de acordo com a função objetivo

Define o menor custo da partícula (ou seja, o custo de $P(i)$) $C^P(i) = C(i)$

Passo 2: Iteração:

Encontra na população a partícula com menor custo e define as variáveis de posição G e custo C^G global com os valores dessa partícula

(*Opcional*) Define o peso de inércia w . Em alguns casos, utiliza-se como critério de parada o número de iterações e faz-se $w = (maxit - iter)/maxit$, onde $maxit$ é o número máximo de iterações e $iter$ o número da iteração atual

Para cada partícula i : Gera números aleatórios r_1 e r_2 segundo uma distribuição uniforme $[0, 1]$ e calcula a nova velocidade da partícula segundo a equação:

$$v(i) = wv(i) + r_1c_1(P(i) - X(i)) + r_2c_2(G - X(i)) \quad (3-14)$$

Calcula a nova posição da partícula segundo a equação:

$$X(i) = X(i) + v(i) \quad (3-15)$$

Avalia o custo da partícula $C(i)$ segundo a função objetiva

Se o custo atual da partícula $C(i)$ for inferior ao menor custo da partícula $C^P(i)$, ou seja, $C(i) < C^P(i)$, define $C^P(i) = C(i)$ e $P(i) = X(i)$

Se o menor custo da partícula $C^P(i)$ for inferior ao menor custo global C^G , ou seja, $C^P(i) < C^G$, define $C^G = C^P(i)$ e $G = X(i)$

Passo 3: Critério de parada:

Avalia o critério de parada (um critério de parada é o número de iterações)

Para o algoritmo se o critério de parada for satisfeito, a melhor solução encontrada é G . Vai para o passo 2 caso contrário

3.4.1 Otimização PSO Inteira

A otimização PSO padrão trabalha com números contínuos. Em algumas aplicações, as soluções para o problema precisam ter valores inteiros. Para esse caso, há uma variação da PSO padrão que discretiza a posição e a velocidade das partículas segundo a equação [49]:

$$\text{INT}(r) = \begin{cases} \text{floor}(r) & \text{se } rand > r - \text{floor}(r) \\ \text{ceil}(r) & \text{caso contrário} \end{cases} \quad (3-16)$$

onde $\text{floor}(r)$ e $\text{ceil}(r)$ são funções de arredondamento para o maior inteiro menor que r e menor inteiro maior que r , respectivamente, e $rand$ é um número aleatório gerado segundo uma distribuição uniforme $[0, 1]$. Utiliza-se neste algoritmo técnica de otimização PSO inteira, pois a solução do problema de otimização na alocação de recursos em redes LTE é um vetor de índices inteiros (índices dos usuários alocados aos blocos de recursos).

3.4.2 Alocação de Recursos Utilizando PSO

O problema de otimização pode ser resolvido através da PSO padrão, porém fazendo algumas considerações. A PSO padrão não possui restrições, então a restrição da banda mínima é convertida em uma função de penalidade. O vetor solução x , que representa o usuário alocado a cada bloco de recurso, é inteiro e é adequado utilizar a versão modificada da PSO, PSO para soluções inteiras, como descrito anteriormente [42].

A função de penalidade descrita pela equação 3-17 converte a otimização com restrições em uma otimização sem restrições [42]:

$$\text{Penalidade} = [\max(R)]^2 \sum_{k=1}^K \left[\min \left(0, \frac{r_k - R_k}{R_k} \right) \right]^2. \quad (3-17)$$

A função de penalidade está associada ao percentual da taxa mínima que foi atendida para cada usuário. Quando as taxas mínimas de todos os usuários são atendidas, a função de penalidade é igual a zero, ou seja, a restrição é atendida [42].

Se o sistema não dispõe de recursos suficientes para atender todas as taxas mínimas, a otimização ainda terá uma solução, mesmo que a restrição não seja plenamente atendida [42].

$$F = T_b - [\max(R)]^2 \sum_{k=1}^K \left[\min \left(0, \frac{r_k - R_k}{R_k} \right) \right]^2, \quad (3-18)$$

$$F = \sum_{k=1}^K r_k(t) - [\max(R)]^2 \sum_{k=1}^K \left[\min \left(0, \frac{r_k - R_k}{R_k} \right) \right]^2. \quad (3-19)$$

A otimização PSO realiza a otimização avaliando os custos de cada solução (partícula) através da função objetivo descrita pela equação 3-19. Os menores custos são memorizados, por partícula e da população, e utilizados no algoritmo PSO [42].

3.5 Alocação de Recursos via Modelagem de Tráfego

Nesta seção é apresentado um esquema de alocação de blocos de recursos em sistemas LTE que visa atender a requisitos de QoS baseando-se na teoria de banda efetiva. Para tal, é apresentado um algoritmo para estimação adaptativa dos parâmetros do modelo de tráfego MWM (*Multifractal Wavelet Model*) com a finalidade de calcular a banda efetiva para os fluxos de tráfego de entrada no sistema [45] [46].

3.5.1 Modelagem Multifractal β MWM Adaptativa

O *Multifractal Wavelet Model* (MWM) é um modelo multifractal com grande destaque na modelagem de tráfego de redes [37] [36]. Ele é baseado em uma cascata multiplicativa no domínio *wavelet*. A transformada *wavelet* [46] [12] discreta é usada neste modelo devido a sua capacidade de representação multiescala de sinais. O MWM apresenta mais de uma modelagem para os coeficientes *wavelet* e coeficientes de escala, gerados pela transformada *wavelet*. Uma dessas modelagem é o β MWM.

O processo de modelagem do β MWM realiza a transformada discreta de *wavelet* de Haar para um número fixo de camadas, J , da cascata multiplicativa binomial [38] para a série completa em uma única etapa. A partir dos coeficientes *wavelet* ($W_{j,i}$) e coeficientes de escala ($U_{j,i}$) gerados, para a camada $j \ \forall \ 0 \leq j \leq J-1$, os parâmetros MWM são estimados [45] [46].

Uma vez estimados os coeficientes de escala $U_{j,i}$, o β MWM assume que os coeficientes de escala da primeira camada ($j = 0$) são independentes e identicamente distribuídos (i.i.d.) e, utilizando o teorema do limite central, possuem uma distribuição normal. Os parâmetros média μ_c e variância σ_c^2 dos coeficientes de escala da primeira camada - agora chamados de $U_{0,0}$ - são estimados para o modelo. Os outros coeficientes de escala podem ser estimados através dos coeficientes *wavelet* e dos coeficientes de escala da primeira camada [45] [46].

Os fatores da cascata multiplicativa são modelados, para a camada j , segundo uma distribuição Beta simétrica [41]. A distribuição de probabilidade Beta possui dois parâmetros. Quando ambos os parâmetros possuem o mesmo valor, a distribuição é chamada de Beta simétrica, pois ela passa a possuir a propriedade de ser simetricamente distribuída entre $[-1, 1]$ e ter média igual a 0. Há variações dessa distribuição onde a

média é $1/2$ e possui algumas propriedades diferenciadas, porém neste trabalho é adotada a distribuição Beta com função densidade de probabilidade dado pela equação 3-20 [41]:

$$f(x) = \frac{(1+x)^{p-1}(1-x)^{p-1}}{\text{Beta}(p, p)2^{2p-1}}, \quad (3-20)$$

onde $\text{Beta}(\cdot, \cdot)$ é a função beta e p é o parâmetro que determina a forma da distribuição.

O β MWM relaciona o decaimento de energia dos coeficientes *wavelet* n_j por camada j com os valores dos parâmetros p_j das distribuições beta simétricas, utilizadas para modelar os multiplicadores da cascata. O decaimento de energia dos coeficientes *wavelet* é dado por [37]:

$$n_j = \frac{E[W_{j-1,k}^2]}{E[W_{j,k}^2]}, \quad (3-21)$$

e os parâmetros p_j são estimados, recursivamente, por:

$$p_j = \frac{n_j}{2} (p_{j-1} + 1) - \frac{1}{2}. \quad (3-22)$$

Os parâmetros do β MWM são estimados adaptativamente. Ao invés do processamento de todos os dados da série de tráfego em uma única etapa, é realizado o processamento iterativo em janelas de tamanho fixo de 2^J amostras, onde J é o número de camadas da cascata. Apenas algumas variáveis são armazenados no processo de modelagem, não havendo a necessidade de guardar uma grande quantidade de dados sobre o fluxo [45] [46].

A modelagem para estimação adaptativa dos parâmetros do β MWM está descrita no Algoritmo 3.3. Assim, com o algoritmo adaptativo obtém-se os parâmetros do modelo β MWM que são: p_j s, μ_c e σ_c^2 [45] [46].

O processo estocástico a partir do modelo β MWM, na camada n , é dado por [37]:

$$C^n[k] = 2^{-n} U_{0,0} \prod_{j=0}^{n-1} (1 + \beta(p_j, p_j)), \quad (3-23)$$

$$C^n[k] = 2^{-n} \text{Norm}(\mu_c, \sigma_c^2) \prod_{j=0}^{n-1} (1 + \beta(p_j, p_j)), \quad (3-24)$$

onde $\beta(\cdot, \cdot)$ é uma variável aleatória beta com p.d.f. dada pela equação 3-20 e $\text{Norm}(\mu, \sigma^2)$

é uma variável aleatória normal com média μ e variância σ^2 .

Algoritmo 3.3: Algoritmo para Estimação Adaptativa dos Parâmetros do β MWM
[45] [46]

início

Inicia as variáveis do modelo

Define o segundo momento dos coeficientes *wavelet*, $E[W_{j,k}^2](0) = 0$

Define a média e variância dos coeficientes de escala $\mu_c(0) = 0$ e $\sigma_c^2(0) = 0$

Define o contador de janela $n = 0$

para Cada nova janela de dados de 2^J amostras **faça**

Realiza a transformada de Haar na janela não sobreposta de dados de 2^J amostras. A transformada de Haar em cada janela de 2^J amostras gera 2^j coeficientes *wavelet* - nomeados de $\tilde{W}_{j,k}$ - por camada j e um coeficiente de escala - nomeado de $\tilde{U}_{0,0}$ - na camada $j = 0$

Atualiza o segundo momento $E[W_{j,k}^2]$ dos coeficientes *wavelet* através da equação:

$$E[W_{j,k}^2](n+1) = E[W_{j,k}^2](n) \left(\frac{n}{n+1} \right) + \frac{\sum_{i=0}^{2^j-1} \tilde{W}_{j,i}^2}{(n+1)2^j} \quad (3-25)$$

Recalcula as taxas de energia n_j segundo a equação (3-21) e os parâmetros p_j segundo a equação (3-22)

Atualiza as estatísticas dos coeficientes de escala segundo as equações:

$$\mu_c(n+1) = \mu_c(n) \left(\frac{n}{n+1} \right) + \frac{\tilde{U}_{0,0}}{n+1} \quad (3-26)$$

$$\begin{aligned} \sigma_c^2(n+1) = & (\sigma_c^2(n) + (\mu_c(n))^2) \left(\frac{n}{n+1} \right) \\ & - (\mu_c(n+1))^2 + \frac{(\tilde{U}_{0,0})^2}{n+1} \end{aligned} \quad (3-27)$$

Incrementa o valor da variável n em 1

3.5.2 Banda Efetiva Utilizando Modelagem β MWM Adaptativa

A banda, no contexto de redes, quantifica a taxa na qual o enlace de rede ou caminho de rede pode transferir dados. A banda efetiva representa a banda necessária para atender requisitos de QoS exigidos para um fluxo, como probabilidade de perda em uma conexão dado um certo tamanho de *buffer* [45] [46].

Seja $X[0, t]$ o tráfego acumulado durante o intervalo de tempo $[0, t]$ para um fluxo de tráfego e que $X[0, t]$ tenha incrementos estacionários, ou seja, $X[0, t + \tau] - X[0, \tau] \stackrel{d}{=} X[0, t] - X[0, 0]$ ($\stackrel{d}{=}$ representa igualdade de distribuição). A banda efetiva do fluxo de tráfego é definida pela equação [29]:

$$\alpha(s, t) = \frac{1}{st} \ln(E[e^{sX[0, t]}]) \quad s > 0, t < \infty. \quad (3-28)$$

Segundo essa definição, a banda efetiva de um processo depende de um parâmetro de espaço s e de um parâmetro de tempo t . Os parâmetros s e t determinam os requisitos de QoS exigidos para o fluxo. A escolha dos parâmetros depende, além dos requisitos de QoS, das características do tráfego [45] [46].

A banda efetiva tem como limite inferior a taxa média ($s \rightarrow 0$) e limite superior a taxa de pico ($s \rightarrow \infty$) do fluxo de tráfego. Quando fluxos de tráfego são simultaneamente servidos a taxas equivalentes as suas bandas efetivas, requisitos de QoS são atendidos para estes fluxos [17].

A partir das equações 3-28 e 3-24 a banda efetiva para o modelo β MWM pode ser escrita na forma [45] [46]:

$$\alpha(s, 2^n) = \frac{1}{st} \ln \left(E \left[e^{\left[2^{-J+n} U_{0,0} \prod_{j=0}^{J-1-n} (1 + \beta(p_j, p_j)) \right]} \right] \right). \quad (3-29)$$

Como a cascata multiplicativa do MWM é diádica, o valor de t deve ser diático, ou seja, $t = 2^n$, $0 \leq n \leq J - 1$.

Um outro método para estimativa da banda efetiva é a estimativa direta. A estimativa direta da banda efetiva é calculada através da média das amostras de tráfego durante o tempo [45] [46].

Sendo N o tamanho da série de tráfego, t_i o tempo de amostra x_i e $I(\tau \leq t_i \leq \tau + t)$ é igual a 1 se t_i está no intervalo $\tau \leq t_i \leq \tau + t$ e 0 caso contrário. A estimativa direta pode ser calculada conforme a seguinte equação [45] [46]:

$$\alpha(s, t) = \frac{1}{st} \ln \left(\frac{1}{N-t} \int_0^{N-t} e^{s \sum_{i=1}^N x_i I(\tau \leq t_i \leq \tau+t)} d\tau \right). \quad (3-30)$$

As Figuras 3.2, 3.3 e 3.4 representam a banda efetiva calculada utilizando modelagem β MWM Adaptativa, conforme equação 3-29, e a banda efetiva calculada de forma direta, conforme equação 3-30, utilizando probabilidade de transbordo do *buffer* de

1%. Pode ser observado que no geral os valores de banda efetiva estimados por meio dos parâmetros adaptativos βMWM tendem a se encontrar com os valores de banda efetiva calculados de forma direta.

Cada um dos 30 usuários citados nas Figuras 3.2, 3.3 e 3.4 é representado em termos de tráfego por meio de diferentes séries reais de tráfego TCP/IP entre a Universidade de Waikato com redes externas, coletados entre 20/05/2011 e 29/10/2011 [48]. Estas séries foram utilizadas a fim de calcular a banda efetiva e tem suas características detalhadas no Apêndice A.

3.5.3 Taxa Mínima Requerida na Alocação de Recursos

O modelo de tráfego βMWM Adaptativo [46] pode ser utilizado para estimação de banda efetiva devido a sua capacidade de cálculo de banda efetiva em tempo real. Uma vez atendido o valor da capacidade indicada pela banda efetiva, o requisito de QoS de probabilidade de perda exigido para o tráfego deverá ser atendido [45] [46].

Seja BE_k a banda efetiva para o usuário k calculada adaptativamente e R'_k a taxa mínima definida para o usuário, define-se a taxa mínima para o usuário k [45] [46]:

$$R_k = \max(BE_k, R'_k) . \quad (3-31)$$

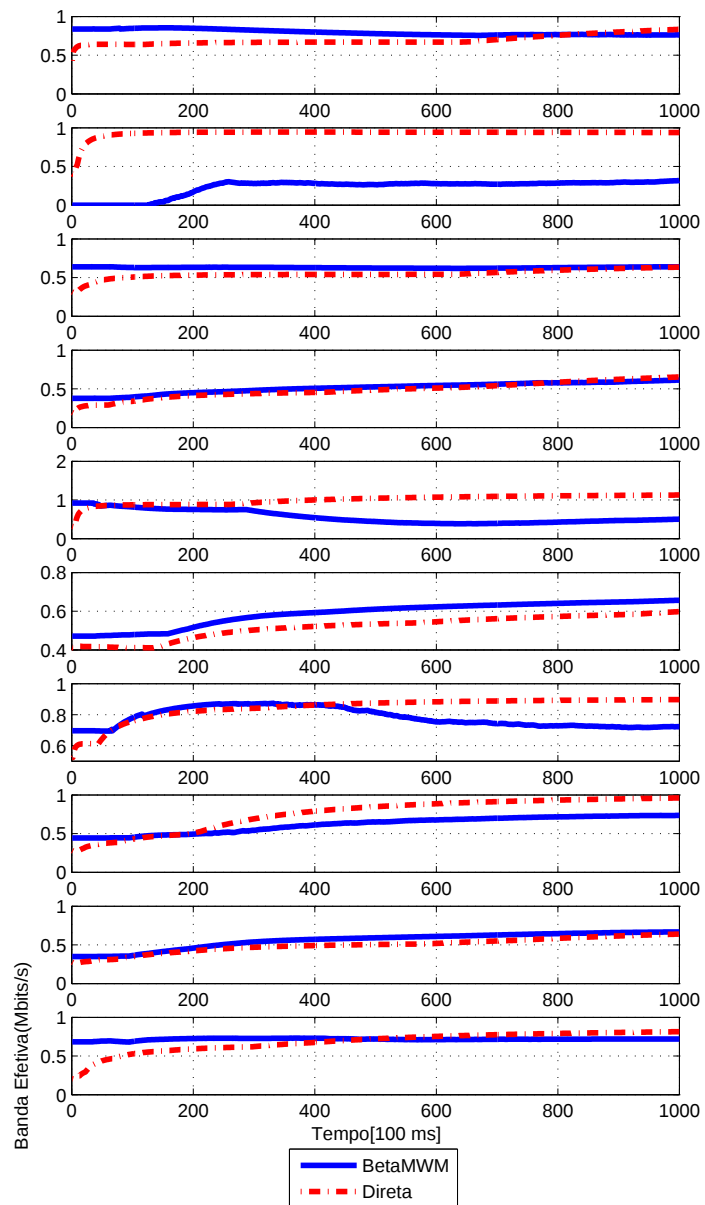


Figura 3.2: Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 1 até 10

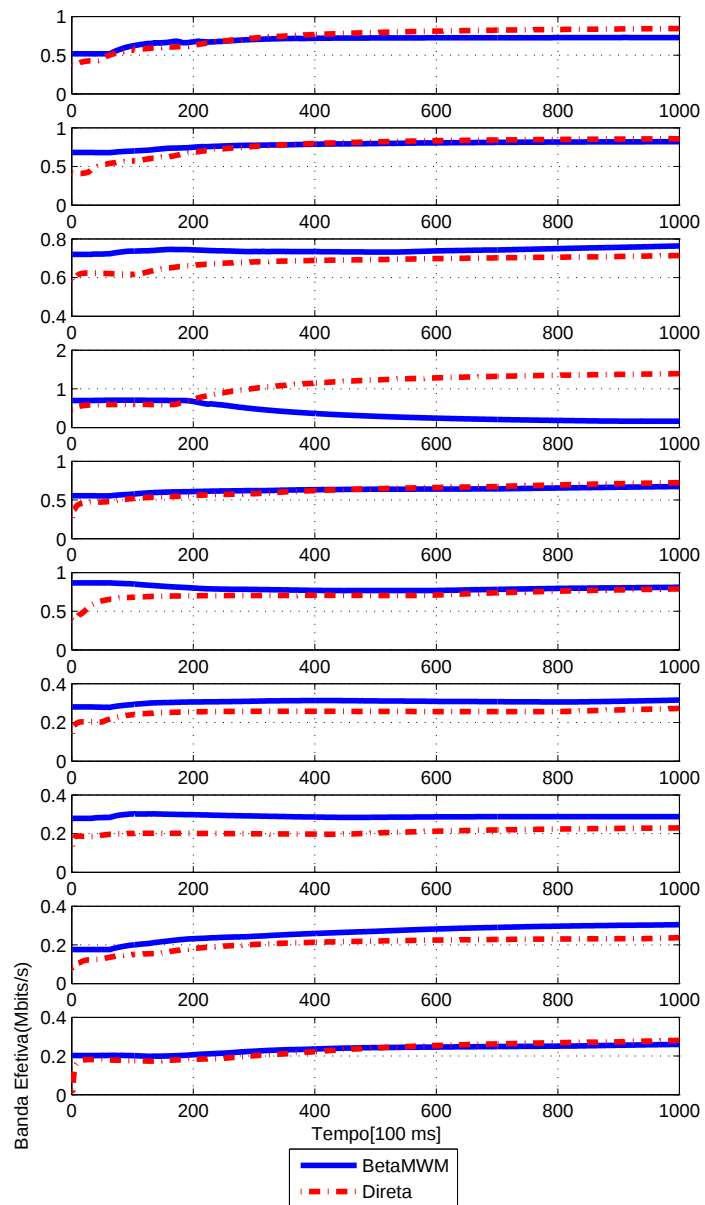


Figura 3.3: Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 11 até 20

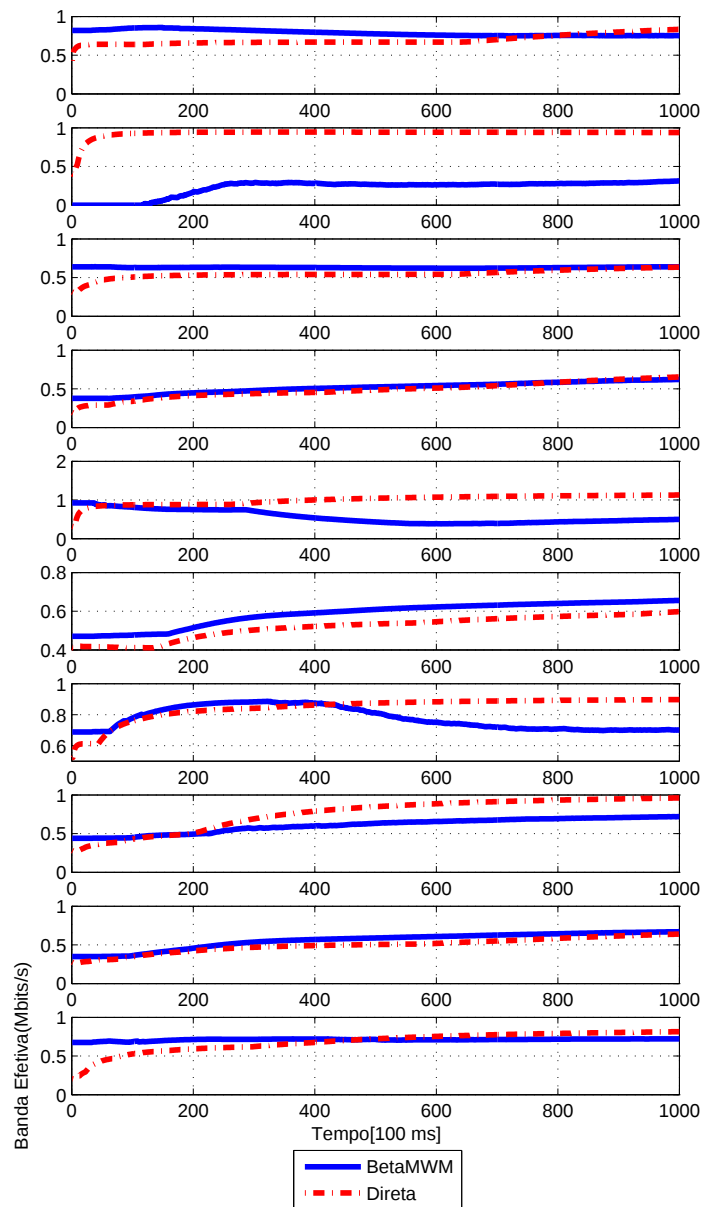


Figura 3.4: Banda efetiva calculada utilizando modelagem β MWM Adaptativa e estimativa direta, para usuários 21 até 30

Redução de Retardo na Alocação de Recursos em Redes LTE

Neste capítulo é proposto um algoritmo de alocação de blocos de recurso para sistemas de comunicação LTE (*Long Term Evolution*) no qual são levadas em conta as restrições impostas pelo esquema MCS (*Modulation and Coding Scheme*) para transmissão *downlink*. O sistema proposto considera a qualidade de transmissão do canal e o critério de retardo máximo para decidir sobre a alocação de recursos de rádio disponíveis, buscando reduzir o retardo médio do sistema. São realizadas comparações com outros algoritmos de alocação através de parâmetros de QoS (*Quality of Service*) como retardo médio, vazão, tempo de processamento, taxa de perda e índice de justiça (*fairness*), comprovando a eficiência do algoritmo proposto [18] [19].

4.1 Algoritmo de Alocação Proposto

O algoritmo PSO (*Particle Swarm Optimization*) apresentado em [42] e o algoritmo com garantia de QoS apresentado em [24] são alguns dentre vários algoritmos de alocação de recursos de rádio para sistemas LTE, porém com critérios e objetivos diferentes. O primeiro algoritmo objetiva maximizar a vazão total do sistema utilizando a heurística de otimização PSO, porém apresenta limitações em relação à complexidade computacional e ao retardo. Enquanto o algoritmo QoS garantido objetiva minimizar a complexidade computacional do problema de otimização da vazão total atendendo a restrição de taxa mínima de transmissão, porém com limitações em termos de taxa de perda.

Propõe-se neste capítulo um algoritmo que tem como objetivo reduzir o retardo médio do sistema, parâmetro de QoS essencial para aplicações em tempo real com taxa de transmissão variável como serviços de VoIP (*Voice over Internet Protocol*) e de videoconferência, que possuem requisitos específicos de banda e retardo máximo tolerável. O algoritmo proposto é uma variação do algoritmo QoS garantido, porém com atenção ao critério de retardo ao invés do critério de taxa mínima de transmissão, e também com critérios de prioridade e de estimativa de blocos diferentes. Ao mesmo

tempo, verifica-se o desempenho em rede do algoritmo proposto a fim de garantir índices de vazão e taxa de perda satisfatórios, tendo o algoritmo PSO como referência.

O algoritmo proposto, descrito em detalhes no Algoritmo 4.1, pode ser resumido em três fases:

1. Estima-se o número de SBs requeridos para cada usuário com prioridade baseada no ganho médio de canal;
2. Alocam-se os SBs para os usuários de acordo com a prioridade, o valor de retardo real, conforme equação 4-1, e o critério de retardo máximo, definido conforme Tabela 4.2;
3. Alocam-se os SBs remanescentes para os usuários com prioridade definida de acordo com o retardo real.

O valor de retardo real d_k para cada usuário k , dada a taxa alcançada r_k e o tamanho do *buffer* B , é calculado como segue:

$$d_k = \frac{B}{r_k}. \quad (4-1)$$

A prioridade na alocação dos recursos é definida em ordem crescente pelo ganho médio do canal por usuário, ou seja, os usuários com piores condições de canal têm maior prioridade. O ganho médio do canal G_k por usuário k é calculado pela seguinte equação:

$$G_k = \frac{1}{N} \sum_{n=1}^N g_{k,n}, \quad (4-2)$$

onde $g_{k,n}$ é o ganho médio do canal para o usuário k no n -ésimo SB e N é o número de SBs disponíveis para *downlink* LTE.

A quantidade N_k de SBs requeridas para cada usuário k é calculada da seguinte forma, com base nas condições do canal:

$$N_k = \text{ceil} \left(\left(\frac{G_k}{G_1 + G_2 + \dots + G_k} \right) * N \right), \quad (4-3)$$

onde G_k é o ganho médio do canal por usuário k , N é o número de SBs disponíveis para *downlink* LTE e $\text{ceil}(x)$ é uma função de arredondamento para o menor inteiro maior que x .

Após calculada a prioridade de alocação, os SBs com maior CQI são alocados de acordo com a quantidade de SBs estimadas para cada usuário. Depois é calculado o valor de retardo real conforme equação 4-1 e verificado se o critério de retardo, definido conforme Tabela 4.2, é satisfeito. Se o critério não for satisfeito, o algoritmo continua alocando SBs buscando satisfazer o critério.

O algoritmo tende a promover melhoria no critério de justiça (*fairness*) e na taxa de perda pela sua característica de balanceamento, uma vez que prioriza os usuários com piores condições do canal com objetivo de satisfazer o critério de retardo, ao mesmo tempo que aloca uma quantidade maior de SBs para os usuários com melhores condições de canal.

Após verificado se o critério de retardo máximo foi satisfeito para todos os usuários, o algoritmo aloca os SBs remanescentes, se houver, priorizando o valor de retardo para cada usuário, ou seja, os usuários com maiores valores de retardo têm maior prioridade. O objetivo é reduzir o retardo médio depois de satisfeito o critério.

Algoritmo 4.1: Algoritmo de alocação proposto com critério de retardo

Entrada: K : número de usuários
 N : número de SBs
 G : CQI dos usuários (matriz $N \times K$)
 R_k : taxa mínima requerida dos usuários (vetor $1 \times K$)
Saída: r_k : vetor de taxas alcançadas por usuário k

```

1  início
2       $W = \{1, 2, \dots, N\}$ 
3       $S_k = \{\}, k \in \{1, 2, \dots, K\}$ 
4       $r_k = \{\}, k \in \{1, 2, \dots, K\}$ 
5       $d_k = \{\}, k \in \{1, 2, \dots, K\}$ 
6      Calcula  $G_k$  conforme equação 4-2
7      Calcula prioridade 1 em ordem crescente das condições do canal (vetor  $G_k$ )
8      Calcula  $N_k$  conforme equação 4-3
9      Define  $d_{max}$  conforme Tabela 4.2
10     para  $k = 1$  até  $K$  de acordo com prioridade 1 faça
11         Aloca SBs com maior CQI para usuário  $k$  no vetor  $S_k$  de acordo com a
            quantidade estimada  $N_k$  e remove o SB alocado de  $W$ 
12         Procura o SB com menor CQI alocado em  $S_k$  e determina o maior MCS
            alcançado para usuário  $k$ 
13         Calcula a taxa  $r_k$  alcançada para usuário  $k$  em um subframe e calcula  $d_k$ 
            conforme equação 4-1
14         enquanto  $d_k \geq d_{max}$  faça
15             Verifica se há SBs disponíveis ( $W \neq 0$ )
16             Aloca SBs com maior CQI para usuário  $k$  no vetor  $S_k$  de acordo com a
                quantidade estimada  $N_k$  e remove o SB alocado de  $W$ 
17             Procura o SB com menor CQI alocado em  $S_k$  e determina o maior MCS
                alcançado para usuário  $k$ 
18             Calcula a taxa  $r_k$  alcançada para usuário  $k$  em um subframe e calcula  $d_k$ 
                conforme equação 4-1
19     se  $W \neq 0$  (Verifica se há SBs disponíveis)
20         então
21             Calcula prioridade 2 em ordem decrescente do vetor  $d_k$ 
22             para  $k = 1$  até  $K$  de acordo com prioridade 2 faça
23                 Verifica se há SBs disponíveis ( $W \neq 0$ )
24                 Aloca SBs com maior CQI para usuário  $k$  no vetor  $S_k$  de acordo com a
                    quantidade estimada  $N_k$  e remove o SB alocado de  $W$ 
25                 Procura o SB com menor CQI alocado em  $S_k$  e determina o maior MCS
                    alcançado para usuário  $k$ 
26                 Calcula a taxa  $r_k$  alcançada para usuário  $k$  em um subframe

```

4.2 Parâmetros de Simulação

Nesta seção são apresentados os parâmetros considerados para modelagem de canal e para o sistema LTE simulado nas seções posteriores.

4.2.1 Modelagem de Canal

As ondas de rádio chegam ao dispositivo móvel com flutuações na fase e amplitude devido a vários fatores, dos quais se pode citar o deslocamento do usuário e a chegada do sinal por diferentes percursos, e consequentemente, com diferentes atrasos de propagação. Nesse contexto, um receptor em determinada localidade pode experimentar um sinal com dezenas de decibéis de diferença em comparação com um receptor similar numa localidade mais próxima à estação base [23].

A atenuação do sinal é proporcional ao quadrado da distância em espaço livre, mas pode variar segundo a quarta ou quinta potência em áreas construídas devido às reflexões e obstáculos. Quando se observa distâncias de poucos quilômetros, a flutuação do sinal ocorre em torno de um valor médio e têm um período mais longo, efeito denominado desvanecimento lento (*slow fading*), causado por movimentos ao longo de distâncias suficientemente grandes para causar alterações bruscas no percurso do sinal. Analisando distâncias de algumas centenas de metros, as flutuações ocorrem mais rapidamente, causadas por pequenos movimentos do transmissor, receptor e objetos circundantes, efeito denominado desvanecimento rápido (*fast fading*). Numa transmissão sem fio, como é o caso de um sistema LTE, o sinal transmitido está exposto aos dois tipos de perturbações [23].

Pelo fato do desvanecimento rápido ser aleatório, suas propriedades estatísticas são usadas para determinar o desempenho do sistema. Para espaços urbanos constituídos densamente por construções, a distribuição de Rayleigh se aproxima bastante como função de probabilidade. Sendo σ o valor eficaz (rms) do sinal recebido, $\frac{r^2}{2}$ a potência instantânea, essa distribuição pode ser expressa por [23]:

$$p(r) = \frac{r}{\sigma^2} e^{-\left(\frac{r^2}{2\sigma^2}\right)}, \text{ para } 0 < r < \infty. \quad (4-4)$$

Em locais onde há pelo menos um percurso direto sem reflexões entre o emissor e receptor, constituindo uma contribuição dominante no sinal recebido, a função de aproximação utilizada será de distribuição de Rice [23]. Seja A a amplitude máxima do sinal dominante, I_0 a função de Bessel modificada do primeiro tipo e ordem zero, σ o desvio padrão da potência local, $\frac{r^2}{2}$ a potência instantânea e σ^2 potência média local do sinal recebido antes da detecção de envoltória (*envelope detection*), a função de probabilidade de Rice é expressa por [23]:

$$p(r) = \frac{r}{\sigma^2} e^{-\left(\frac{r^2+A^2}{2\sigma^2}\right)} I_0\left(\frac{Ar}{\sigma^2}\right), \text{ para } A \geq 0, r \geq 0. \quad (4-5)$$

Essa distribuição é frequentemente descrita em termos do fator K , conhecido como Fator Rician, que é obtido de [23]:

$$K = 10 \log \left(\frac{A^2}{2\sigma^2} \right) \text{ dB}. \quad (4-6)$$

Se $A \rightarrow 0, K \rightarrow \infty$ dB, a amplitude do caminho dominante decai e a distribuição tende para uma distribuição Rayleigh [23].

Nesse ambiente o comportamento varia também de maneira randômica entre pequenas áreas vizinhas, o que se denomina efeito de sombreamento (*shadow effect*), ou desvanecimento lognormal. Esse componente *slow fading* do meio de transmissão é representada geralmente por uma distribuição lognormal. Dado o sinal, $s(t)$, seu valor médio quadrado ou potência local média é lognormal em dBm com variância igual a σ_s^2 . Seja S_m o valor médio de S em dBm, σ_S o desvio padrão de S em dB, S igual a $10 \log s$ em dBm e s a potência do sinal em Watts, a distribuição é dada por [23]:

$$p(S) = \frac{1}{\sqrt{2\pi}\sigma_S} e^{-\left[\frac{S-S_m}{2\sigma_S^2}\right]}. \quad (4-7)$$

As condições do canal para cada usuário e SB em termos de SINR (*Signal-to-Interference-plus-Noise-Ratio*) foram gerados para cada TTI através da seguinte equação [33]:

$$SINR = \frac{G_o P_o}{\sum_{j=1}^J G_j P_j + \sigma_n^2}, \quad (4-8)$$

onde G_o é o ganho do canal para a potência de transmissão P_o , G_j é o ganho do canal para os sinais de interferência com potência P_j , σ_n^2 é a potência do ruído e J é o número de células de interferência.

A Tabela 4.1 mostra os parâmetros de modelagem de canal considerados em simulação.

Tabela 4.1: *Parâmetros de modelagem de canal*

Modelo multipercurso	Rayleigh / Rician [33]
Perfil de atraso multipercurso	Pedestre Estendido A (EPA), Veicular Estendido A (EVA), Urbano Típico Estendido (ETU) [1]
Modelo de perda de percurso	$L = 128.1 * 37.6 \log_{10}(R)$, R em km [3]
Distância entre UE (<i>User Equipment</i>) e BS (<i>Base Station</i>)	0.7km / 1km / 1.4km
Densidade de potência do ruído branco	-174dBm/Hz [3]
Potência máxima do transmissor BS	46dBm [3]
Ganho da antena BS após perda no cabo	15dBi [3]
Ganho da antena UE	0dBi [3]
Figura de ruído UE	9dB [3]
Margem de interferência UE	4dB [3]
Velocidade UE	3km/h [33]

4.2.2 Parâmetros do Sistema LTE

A Tabela 4.2 mostra os valores de critério de retardo máximo adotados em simulação, considerando cenários com 2 até 30 usuários simultaneamente. Os valores adotados são aleatórios e poderiam ser diferentes.

Tabela 4.2: *Valores de critério de retardo máximo adotados*

Número de usuários (5MHz)	Número de usuários (10MHz)	Retardo (ms)
2 a 4	2 a 8	180
5 a 8	9 a 16	360
9 a 12	17 a 24	540
mais de 13	mais de 25	720

A Tabela 4.3 mostra os parâmetros de transmissão *downlink* do sistema LTE. Inicialmente, são considerados 1000 TTIs, que correspondem a 1 segundo de simulação, e, posteriormente, 864×10^5 TTIs, que correspondem a um período de 24 horas de simulação, de forma a verificar se os resultados diferem para um intervalo de tempo maior. São considerados cenários com 24 e 50 blocos de recursos, correspondentes às larguras de banda de 5MHz e 10MHz, respectivamente. O tamanho do *buffer* por usuário é fixado em 60kB, valor este comumente utilizado na literatura [44]. A taxa de perda na estimativa de banda efetiva, ou seja, a probabilidade de transbordo do *buffer* de 60 kB, é definida em 1%.

A taxa mínima requerida é fixada inicialmente com valores de 0.768, 1.024 e 1.536 Mbits/s, aleatoriamente definidos para cada usuário. Também são considerados como taxa mínima requerida os valores de banda efetiva calculados adaptativamente por modelagem de tráfego, conforme explicado na Seção 3.5. O modelo de tráfego β MWM Adaptativo [46] é utilizado para estimação de banda efetiva devido a sua capacidade de cálculo de banda efetiva em tempo real. Uma vez atendido o valor da capacidade indicada pela banda efetiva, o requisito de QoS de probabilidade de perda exigido para o tráfego deverá ser atendido [45] [46].

Tabela 4.3: *Parâmetros de simulação do sistema LTE*

Número de TTIs simulados	1000 (1 segundo) / 864×10^5 (24 horas)
Número de blocos de recursos	25 / 50
Largura de banda	5MHz / 10MHz
Número de subportadoras por bloco de recurso	12
Largura de banda por subportadora	15kHz
Número de símbolos OFDM por <i>slot</i>	7 (prefixo cíclico normal)
Tempo de escalonamento (TTI)	1ms
Tamanho de <i>buffer</i> por usuário	60kB [44]
Taxa mínima requerida (R_k)	Fixa (0.768; 1.024; 1.536 Mbits/s) [4] / Calculada adaptativamente conforme Seção 3.5 [45] [46] [5]
Taxa de perda na estimativa de banda efetiva	1% [45] [46] [5]

A Tabela 4.4 mostra a taxa de bits e SINRs associados ao MCS. Os valores apresentados nesta tabela são característicos do LTE.

Tabela 4.4: *Taxas de bits e SINRs associadas aos MCS [26]*

Nível MCS	SINR (dB)	MCS	Taxa de bits (kbps)
1	1.7	QPSK (1/2)	168
2	3.7	QPSK (2/3)	224
3	4.5	QPSK (3/4)	252
4	7.2	16QAM (1/2)	336
5	9.5	16QAM (2/3)	448
6	10.7	16QAM (3/4)	504
7	14.8	64QAM (2/3)	672
8	16.1	64QAM (3/4)	756

4.3 Simulações e Resultados

Considerando os parâmetros de transmissão *downlink* no sistema LTE e do modelo de canal descritos na Seção 4.2, avalia-se nesta seção os resultados das simulações do algoritmo de alocação proposto, comparando com o algoritmo PSO (*Particle Swarm Optimization*) [42] e o algoritmo com garantia de QoS [24]. Os valores dos parâmetros de QoS apresentados representam valores médios para os TTIs simulados.

As simulações foram realizadas por meio do software MATLAB versão R2014a, utilizando um microcomputador com a seguinte configuração: Processador Intel Core I5-3570 3.40 GHz, 8GB RAM, HD SATA III 7200 RPM, Windows 8 64bits. Optou-se por implementar as funções e rotinas de simulação ao invés de utilizar ferramentas disponíveis de simulação de rede a fim de simular diversos parâmetros de modelagem de canal e do sistema LTE.

4.3.1 Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 1000 TTIs

Nesta subseção considerou-se distância entre estação base e usuário de 0.7km, largura de banda de 5MHz, taxa mínima fixa, 1000 TTIs e os canais multipercurso Rayleigh e Rician, utilizando fator K para distribuição Rician com valores iguais a 10 e 100. São considerados cenários com 2 até 20 usuários em simulação.

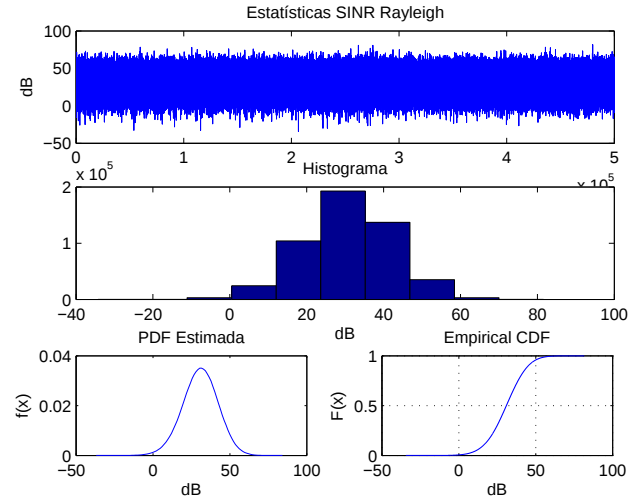
A Tabela 4.5 apresenta dados estatísticos referentes à modelagem de canal simulada em termos de SINR. O desvio padrão, a variância, o índice de dispersão e a variância normalizada são medidas comuns de dispersão estatística que medem o quanto de variação dos dados existe em relação a média. O momento de 3ª ordem mede a simetria dos dados em torno do valor médio, enquanto o momento de 4ª ordem mede a concentração (acuidade) dos dados em torno do valor médio.

Nota-se pela Tabela 4.5 que as três modelagens apresentam dados estatísticos similares. Isto significa que as modelagens são semelhantes em termos de SINR. Optou-se por utilizar diferentes modelos multipercurso a fim de avaliar o desempenho do algoritmo de escalonamento proposto para diferentes cenários.

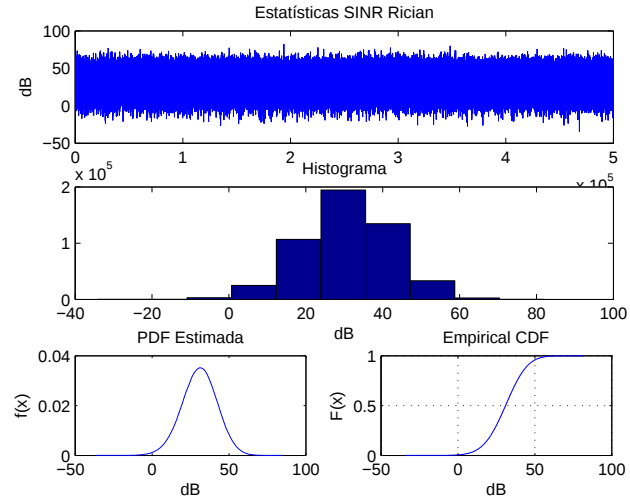
Tabela 4.5: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs

Estatística / Modelo	Rayleigh	Rician (Fator $K = 10$)	Rician (Fator $K = 100$)
Média	30.69	30.74	30.77
Desvio padrão	11.48	11.42	11.43
Variância	131.92	130.61	130.78
Valor mínimo	-34.28	-34.12	-43.43
Valor máximo	81.60	81.99	82.18
Índice de dispersão	4.29	4.24	4.25
Variância normalizada	0.14	0.13	0.14
Momento de 3ª ordem (simetria)	-211.11	-193.78	-201.67
Momento de 4ª ordem (concentração)	54638.5	53317.4	53652.6

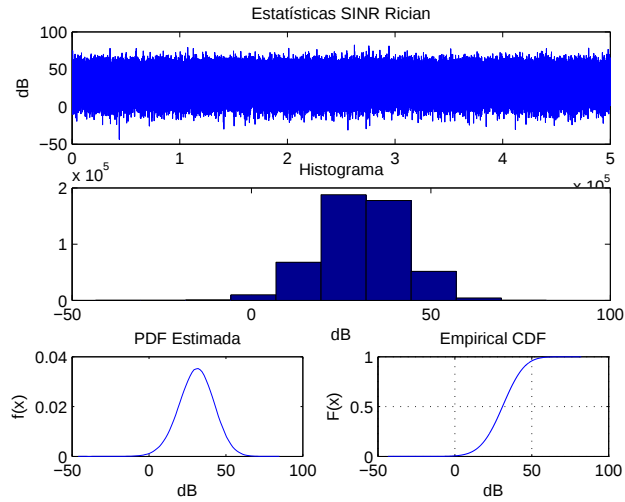
A Figura 4.1 apresenta o histograma, a função densidade de probabilidade (*Probability Density Function*, PDF) e a função de distribuição cumulativa (*Cumulative Distribution Function*, CDF), referentes à modelagem de canal simulada em termos de SINR. Nota-se que as modelagens de canal apresentam distribuições semelhantes.



(a)



(b)



(c)

Figura 4.1: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs

A Tabela 4.6 apresenta a taxa de atendimento ao critério de retardo máximo, definido conforme Tabela 4.2. O algoritmo proposto apresenta maior taxa de atendimento, com valores próximos a 100%, independente do modelo multipercurso considerado. Isto significa que quase 100% dos usuários apresentam valores de retardo menores que o critério máximo definido. O algoritmo QoS garantido é aquele que apresenta a menor taxa de atendimento.

O algoritmo PSO é aquele que apresenta maiores valores para vazão total no geral, conforme ilustrado na Figura 4.2. Este fato era esperado visto se tratar de heurística de otimização da vazão total, em detrimento do tempo de processamento e índice de justiça. O algoritmo proposto apresenta valores de vazão total semelhantes aos apresentados pelo algoritmo QoS garantido, e até superiores aos valores apresentados pelo algoritmo PSO se considerado mais de 14 usuários em simulação. O algoritmo proposto tende a apresentar desempenho superior ao se considerar um número maior de usuários em simulação, devido ao fato de estimar um número maior de blocos para os usuários com melhores condições de canal.

Quanto à vazão média por usuário, ilustrada na Figura 4.3, nota-se que o algoritmo PSO apresenta os maiores valores. O algoritmo proposto e o algoritmo QoS garantido apresentam valores semelhantes.

O algoritmo proposto apresenta no geral os menores valores para o retardo médio por usuário, conforme pode ser visto na Figura 4.4. Os pontos não demarcados nos gráficos significam que o retardo médio tende ao infinito, ou seja, a taxa de transmissão para um ou mais usuários é igual a zero e não há transmissão de dados. Verifica-se que o algoritmo proposto garante transmissão para todos os usuários no cenário simulado, mesmo aqueles com condições degradadas de sinal.

Em relação ao índice de justiça (*fairness*), ilustrado na Figura 4.5, o algoritmo proposto apresenta valores acima de 0.8, tendendo a manter valores mais constantes para diferentes números de usuários considerados em simulação. Tal fato se justifica pela característica de balanceamento do algoritmo proposto, qual seja, os usuários com piores condições de canal têm maior prioridade enquanto os usuários com melhores condições de canal têm um número maior de blocos alocados. Estes valores apresentados pelo algoritmo proposto são no geral superiores aos apresentados pelo algoritmo PSO e, para mais de 12 usuários, aos apresentados pelo algoritmo QoS garantido. O algoritmo QoS garantido tende a apresentar valores de índice de justiça cada vez menores à medida que o número de usuários aumenta, em virtude de sua característica de procurar atender ao critério de taxa mínima de transmissão, o que em alguns casos pode fazer com que os usuários com menor prioridade não sejam atendidos durante a transmissão devido ao número insuficiente de blocos.

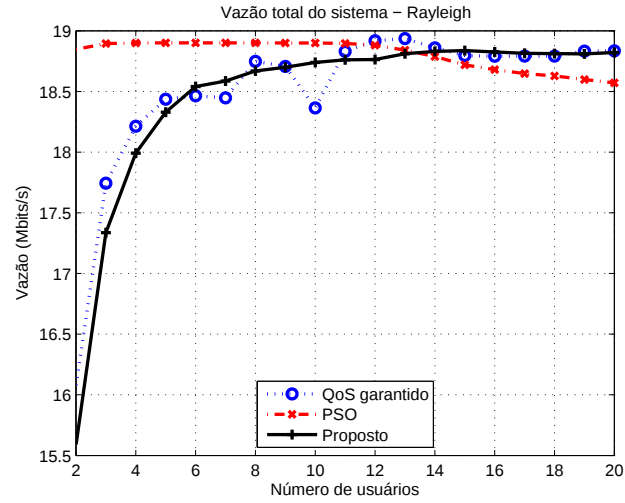
Na Figura 4.6 observa-se que o algoritmo PSO apresenta menores valores para

perda média do sistema dentre todos os algoritmos estudados. Com relação a este parâmetro de QoS, o algoritmo proposto apresenta melhor desempenho em comparação ao algoritmo QoS garantido, com valores de perda média bastante inferiores. A característica de balanceamento do algoritmo proposto explica tal fato, enquanto o algoritmo QoS prioriza o atendimento da taxa mínima.

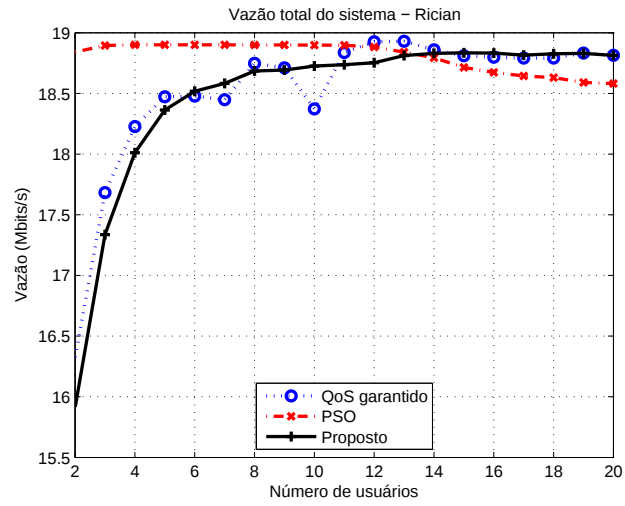
O algoritmo proposto e o algoritmo QoS garantido apresentam valores semelhantes de tempo de processamento médio, conforme pode ser verificado na Tabela 4.7. O algoritmo PSO apresenta os maiores valores de tempo de processamento para os diferentes cenários, o que era esperado por se tratar de heurística de otimização com critérios de parada definidos.

Tabela 4.6: Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs

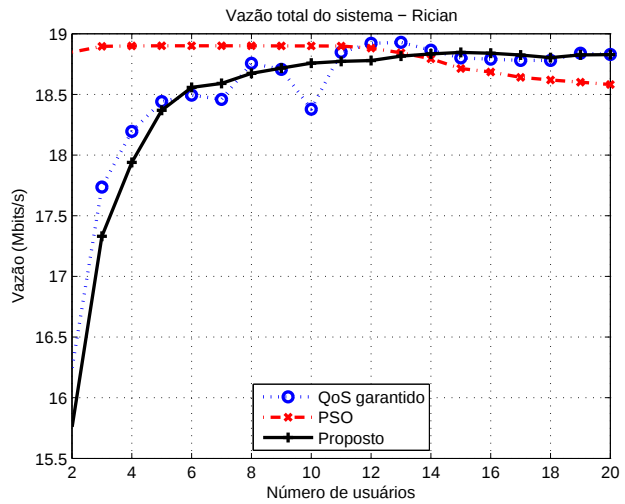
Algoritmo	Modelo multipercurso	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	Rayleigh	53911	74.20
	Rician (Fator $K = 10$)	53974	74.17
	Rician (Fator $K = 100$)	53899	74.21
PSO	Rayleigh	12340	94.09
	Rician (Fator $K = 10$)	12159	94.18
	Rician (Fator $K = 100$)	12205	94.16
Proposto	Rayleigh	1359	99.35
	Rician (Fator $K = 10$)	1375	99.34
	Rician (Fator $K = 100$)	1288	99.38



(a)

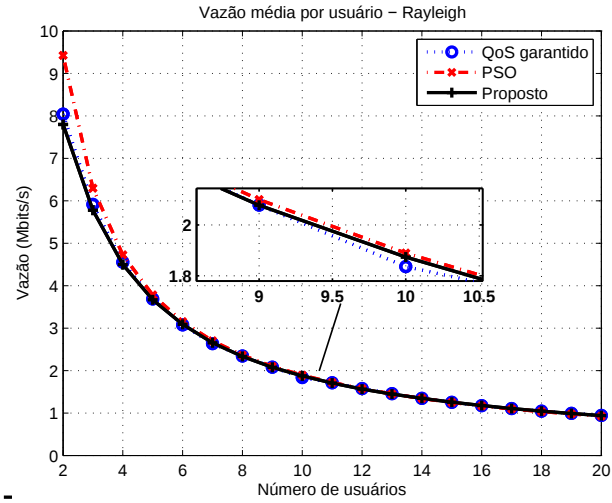


(b)

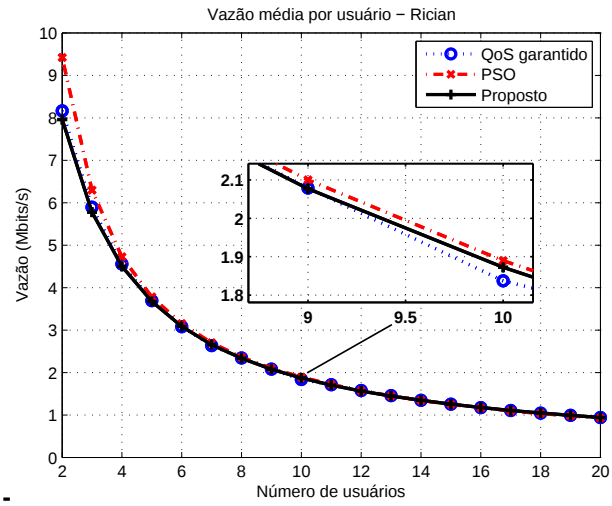


(c)

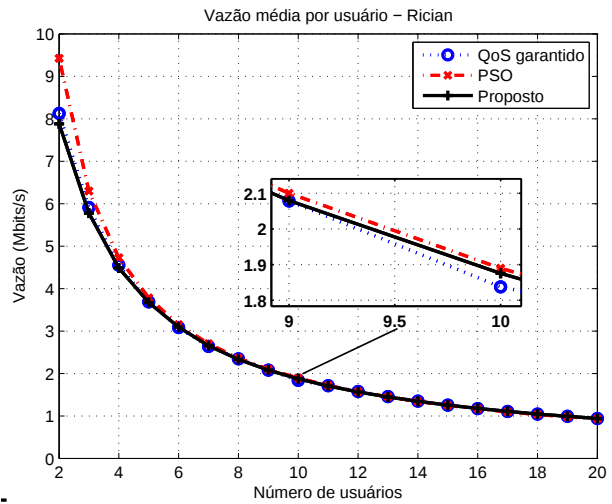
Figura 4.2: Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)

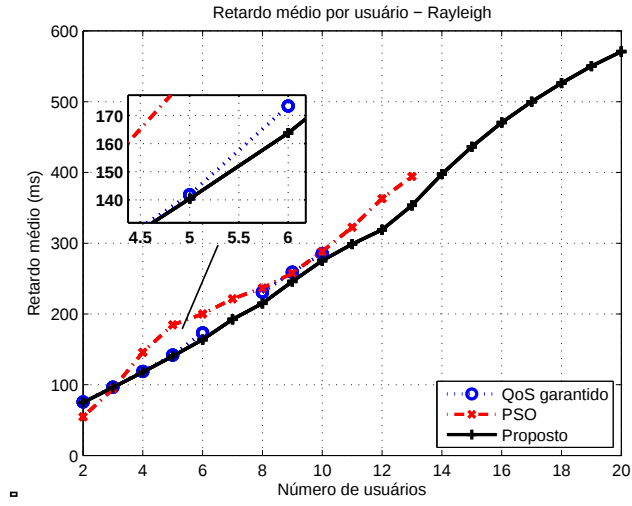


(b)

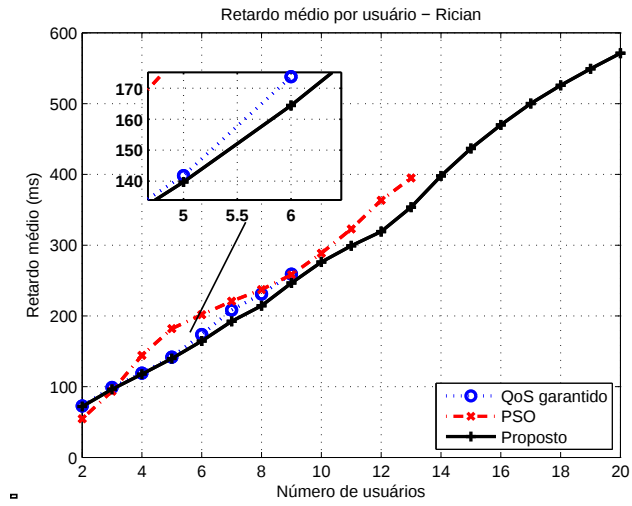


(c)

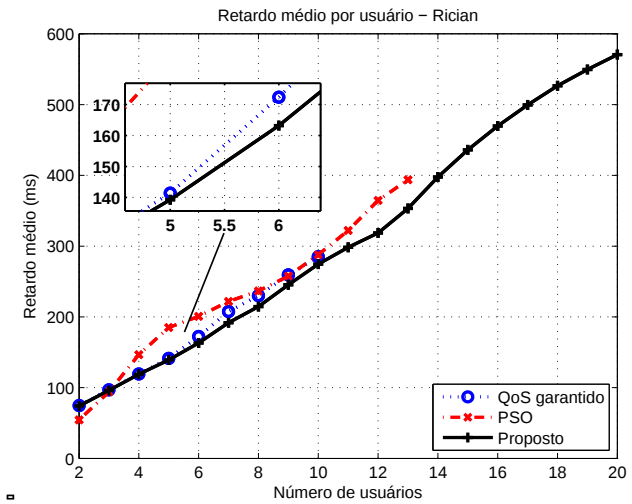
Figura 4.3: Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)

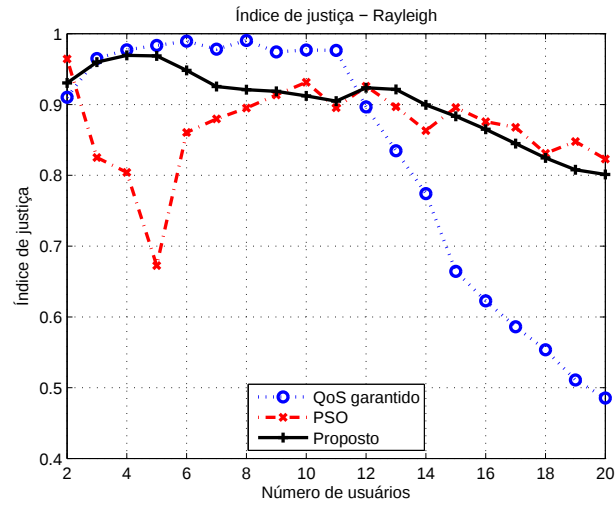


(b)

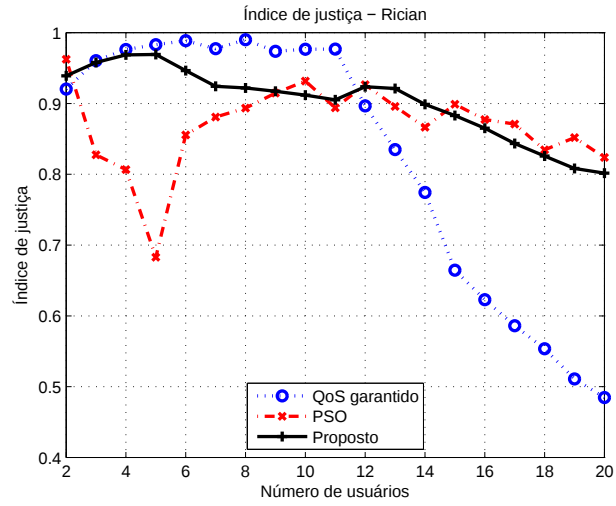


(c)

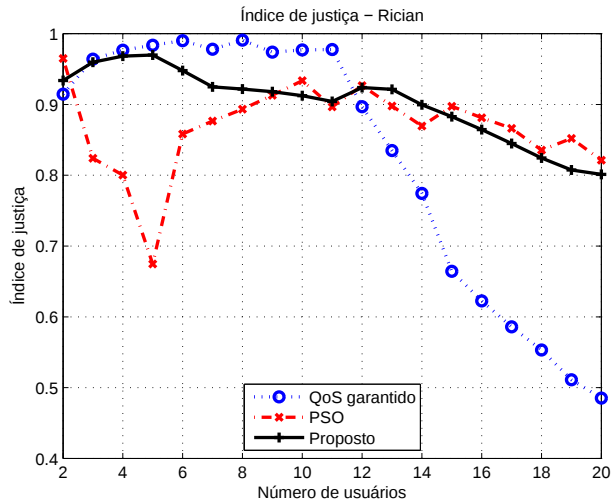
Figura 4.4: Retardo médio entre usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)

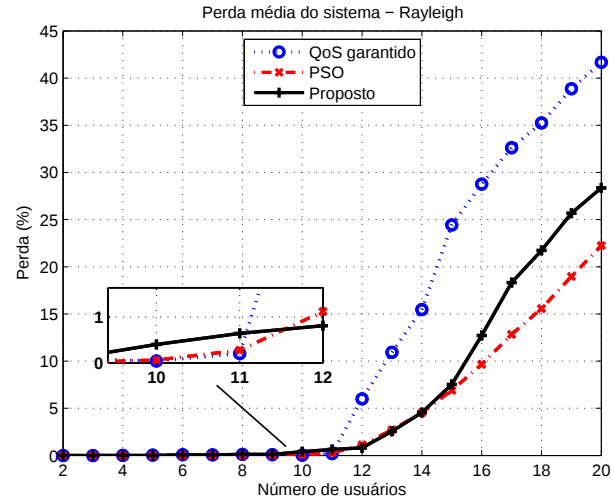


(b)

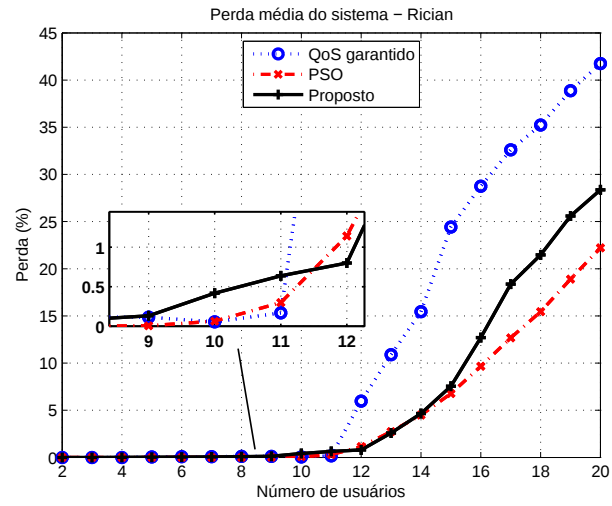


(c)

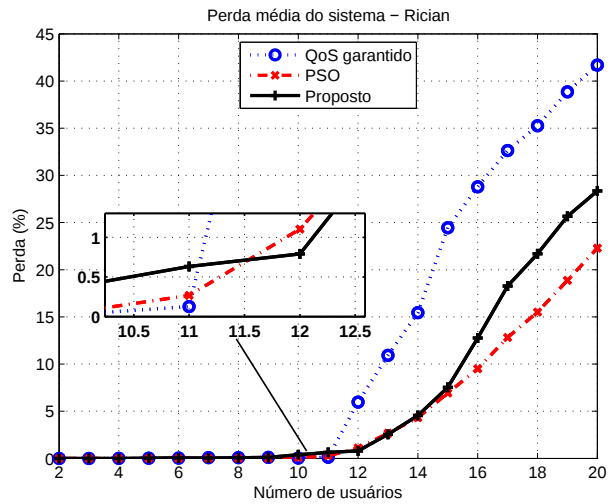
Figura 4.5: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)



(b)



(c)

Figura 4.6: Perda média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs

Tabela 4.7: *Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 5MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs*

Algoritmo	Modelo	Tempo de processamento (ms)				
		4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS	Rayleigh	0.246	0.349	0.448	0.554	0.609
	Rician (Fator $K = 10$)	0.247	0.351	0.452	0.557	0.609
	Rician (Fator $K = 100$)	0.248	0.354	0.448	0.557	0.604
PSO	Rayleigh	69.808	71.101	72.603	73.845	75.578
	Rician (Fator $K = 10$)	69.328	70.869	72.272	73.462	74.763
	Rician (Fator $K = 100$)	70.401	71.723	72.909	74.343	75.456
Proposto	Rayleigh	0.274	0.399	0.506	0.535	0.749
	Rician (Fator $K = 10$)	0.270	0.399	0.506	0.538	0.745
	Rician (Fator $K = 100$)	0.293	0.420	0.529	0.556	0.766

4.3.2 Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 24 Horas de Simulação

Nesta subseção considerou-se uma distância entre usuário e estação base de 0.7km, largura de banda de 5MHz, taxa mínima fixa e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 20 usuários em 864×10^5 TTIs, correspondente a um período de 24 horas.

Os resultados apresentados nesta subseção demonstram que o cenário com 864×10^5 TTIs se assemelha ao cenário apresentado na subseção anterior, com 1000 TTIs.

As Tabelas 4.8 e 4.9 apresentam valores de taxa de atendimento e tempo de processamento similares aos valores apresentados nas Tabelas 4.6 e 4.7. O algoritmo proposto apresenta maior taxa de atendimento ao critério de retardo máximo em relação aos algoritmos PSO e QoS garantido, e apresenta valores similares de tempo de processamento aos valores apresentados pelo algoritmo QoS garantido.

O fato se repete em relação às Figuras 4.7, 4.8, 4.9, 4.10 e 4.11, que mostram gráficos semelhantes aos gráficos apresentados nas Figuras 4.2, 4.3, 4.4, 4.5 e 4.6. O cenário se mantém inalterado, mesmo com um estudo de caso que abrange um tempo de simulação maior.

Tabela 4.8: *Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação*

Algoritmo	Taxa (%)
QoS garantido	74.56
PSO	94.29
Proposto	99.69

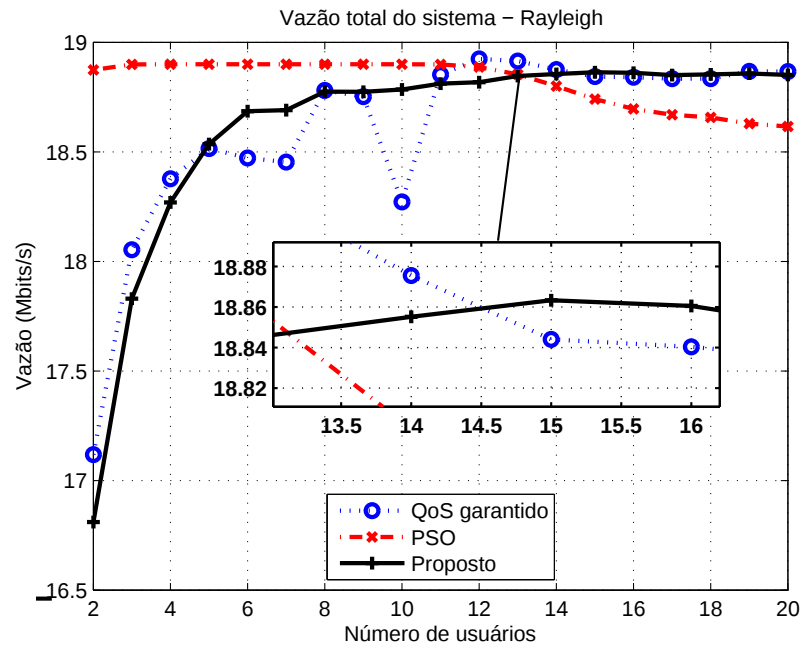


Figura 4.7: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

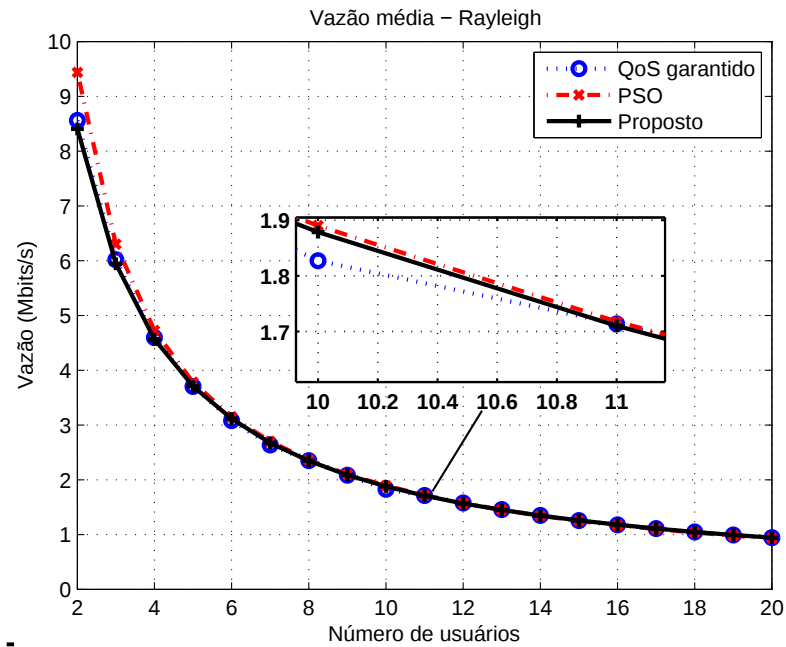


Figura 4.8: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

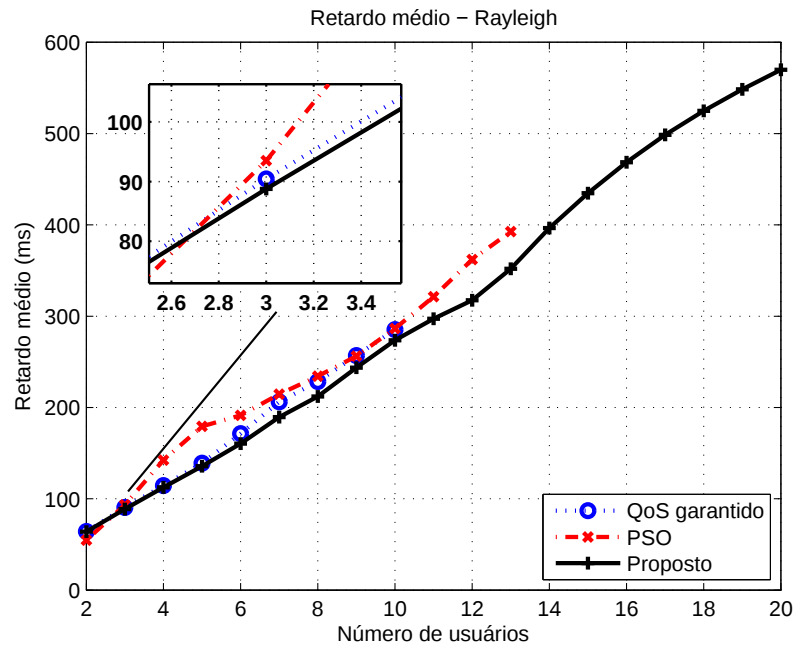


Figura 4.9: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

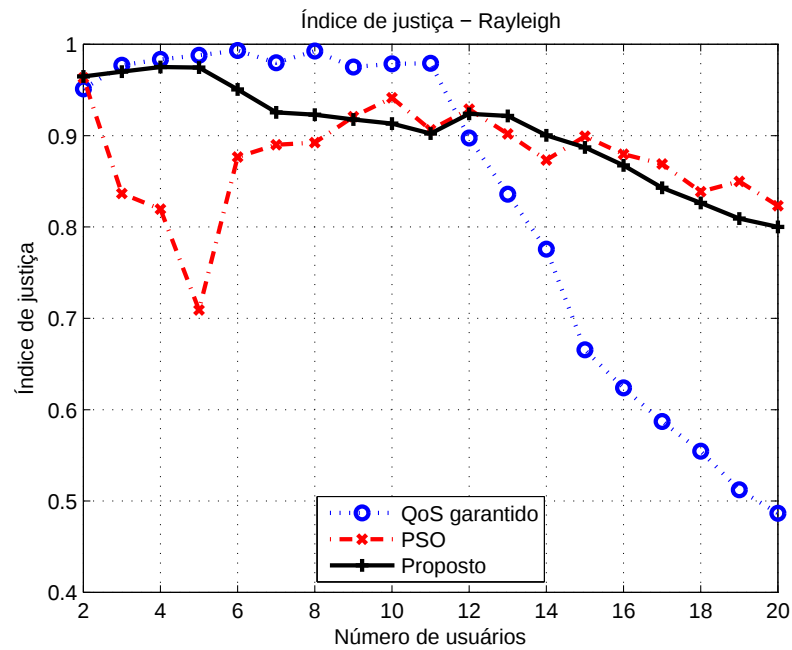


Figura 4.10: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

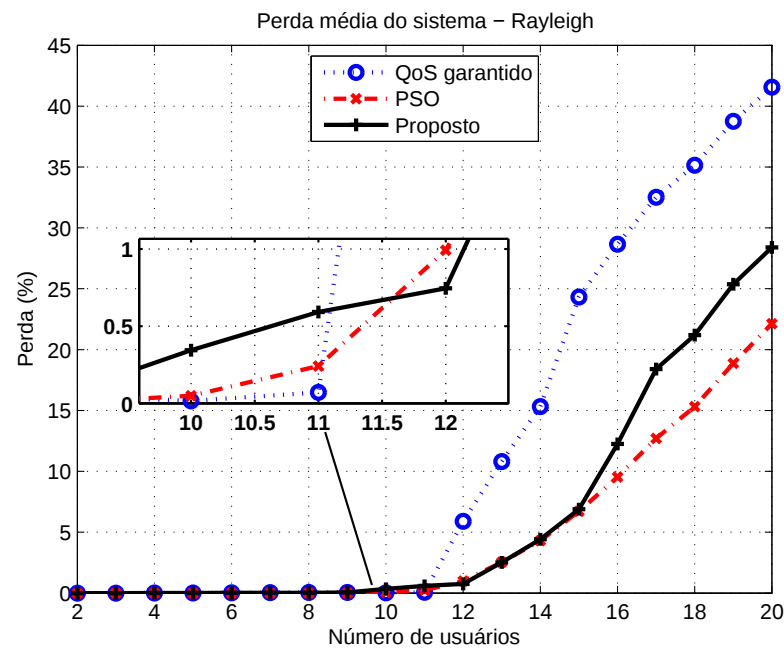


Figura 4.11: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

Tabela 4.9: Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

Algoritmo	Tempo de processamento (ms)				
	4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS garantido	0.254	0.359	0.464	0.576	0.624
PSO	73.191	74.564	75.949	78.169	78.907
Proposto	0.277	0.411	0.520	0.552	0.764

4.3.3 Distância de 0.7km, Largura de Banda de 5MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs

Nesta subseção considerou-se uma distância entre usuário e estação base de 0.7km, largura de banda de 5MHz, taxa mínima calculada adaptativamente conforme explicado na Seção 3.5 e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 20 usuários em 1000 TTIs.

Nota-se na Tabela 4.10 que o algoritmo proposto apresenta uma taxa de atendimento ao critério de retardo de 99.13%, maior que a taxa apresentada pelos algoritmos QoS garantido e PSO.

A Figura 4.12 mostra que o algoritmo PSO apresenta a maior vazão total e inclusive melhora o próprio desempenho em relação ao estudo de caso com taxa mínima fixa, devido ao fato dos valores calculados adaptativamente serem em média menores do que os valores fixados, promovendo consequentemente penalidade menor e portanto maiores valores de vazão total. O algoritmo proposto apresenta vazão total similar à apresentada pelo algoritmo QoS para mais de 8 usuários, em virtude da característica do algoritmo proposto de estimar um número maior de blocos para usuários com melhores condições de canal. Tal fato se repete em relação à vazão média, conforme pode ser visto na Figura 4.13.

Em relação ao retardo médio entre os usuários, observa-se que o algoritmo proposto apresenta os menores valores, conforme verificado na Figura 4.14. O algoritmo PSO apresenta os maiores valores de retardo.

Quanto ao índice de justiça, ilustrado na Figura 4.15, o algoritmo QoS garantido apresenta o melhor desempenho, enquanto o algoritmo PSO apresenta o pior desempenho. Como os valores de taxa mínima calculados adaptativamente são em geral menores que os valores fixos, o algoritmo QoS garantido atende com maior facilidade este critério, e portanto um maior número de usuários são atendidos.

A Figura 4.16 mostra que o algoritmo proposto apresenta a menor taxa de perda média do sistema, enquanto o algoritmo QoS garantido apresenta em geral a maior taxa. O fato se justifica pela característica de balanceamento do algoritmo proposto, que tende a promover melhoria no desempenho da taxa de perda, uma vez que garante nos cenários simulados a transmissão de dados para todos os usuários, independente das condições de canal.

O algoritmo proposto e o algoritmo QoS garantido apresentam tempo de processamento similar, conforme mostrado na Tabela 4.11. O algoritmo PSO apresenta o maior tempo de processamento.

Tabela 4.10: Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	14741	92.95
PSO	24415	88.32
Proposto	1810	99.13

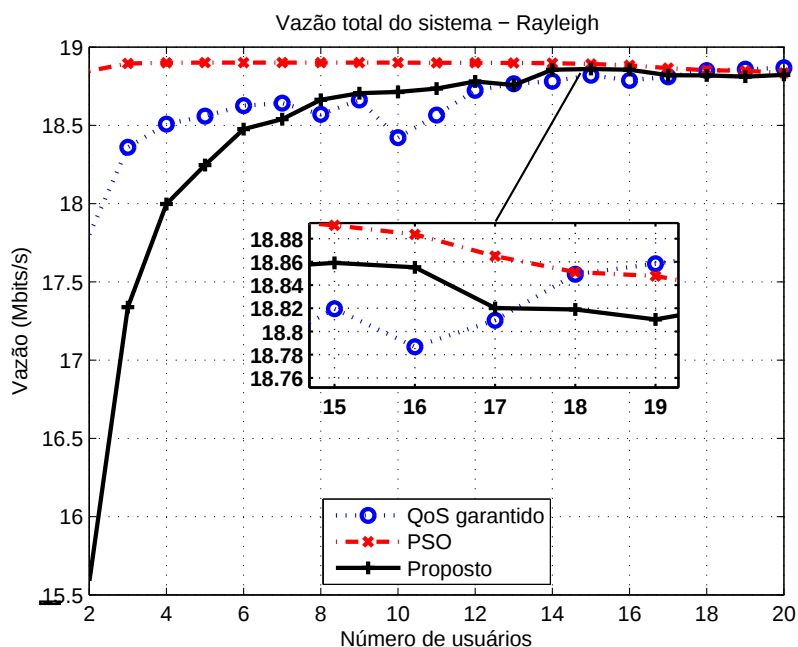


Figura 4.12: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

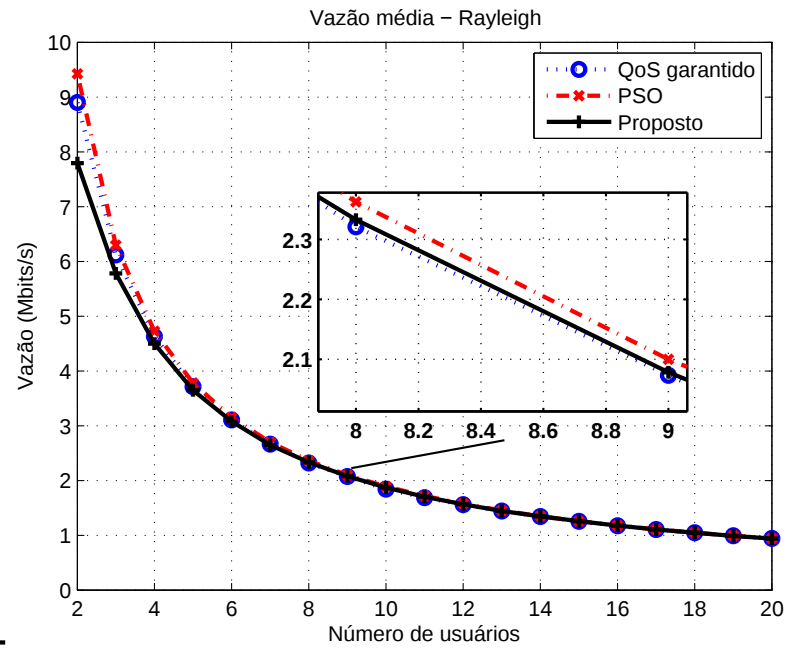


Figura 4.13: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

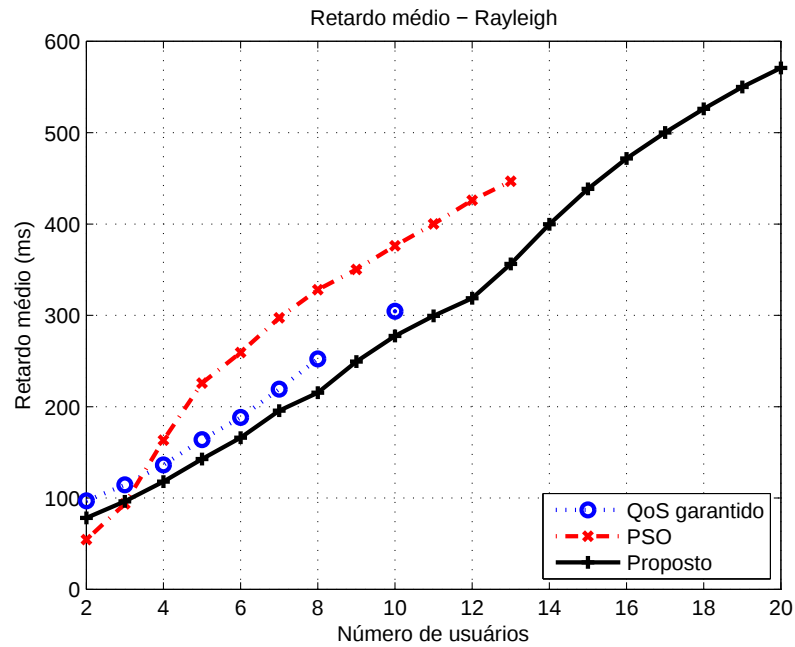


Figura 4.14: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

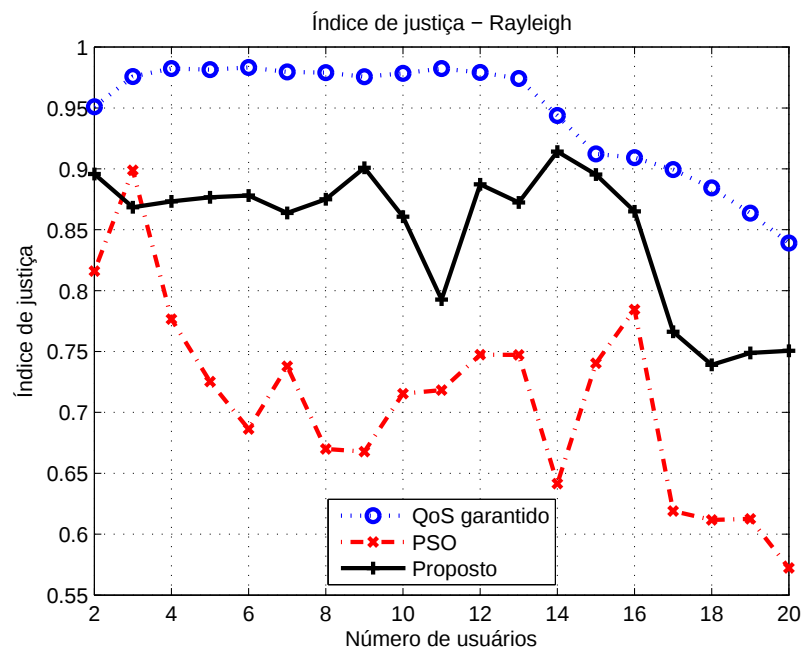


Figura 4.15: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

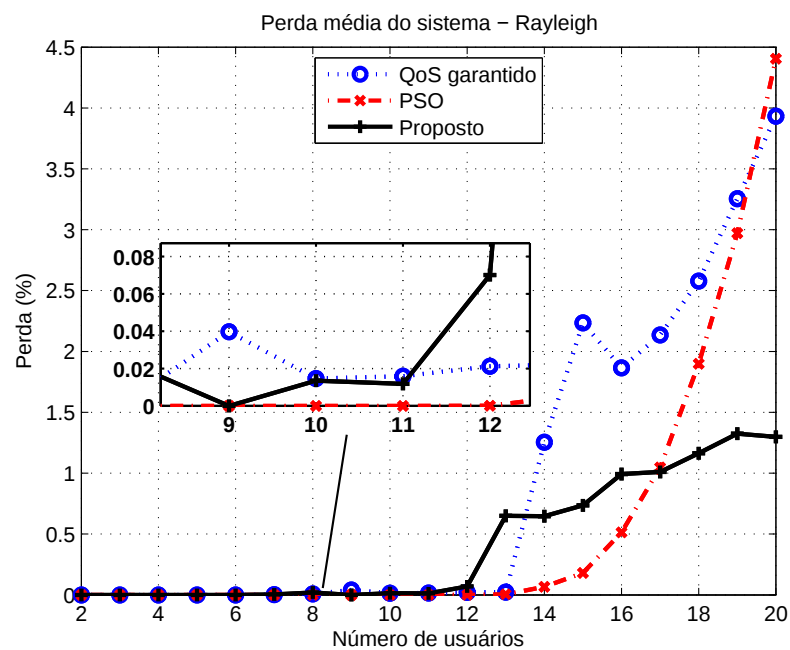


Figura 4.16: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 4.11: *Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs*

Algoritmo	Tempo de processamento (ms)				
	4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS garantido	0.256	0.353	0.448	0.544	0.630
PSO	72.543	74.043	75.412	76.641	78.129
Proposto	0.277	0.400	0.513	0.542	0.752

4.3.4 Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 1000 TTIs

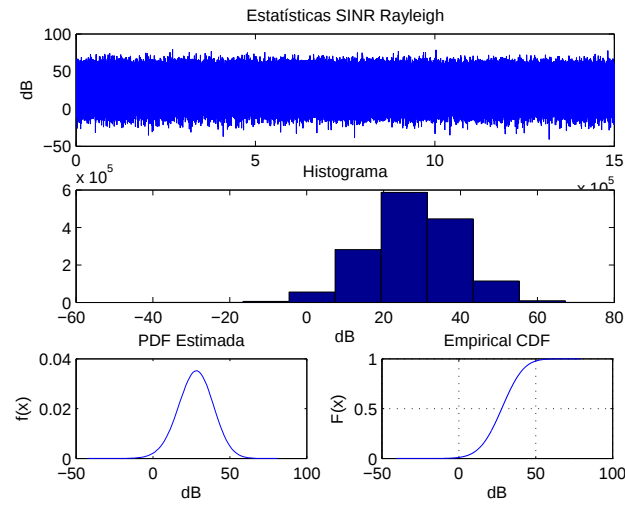
Nesta subseção considerou-se cenário com distância entre usuário e estação base de 0.7km, largura de banda de 10MHz, taxa mínima fixa, 1000 TTIs e canais multipercurso Rayleigh e Rician, com fator K igual a 10 e 100.

Observa-se pela Tabela 4.12 que as modelagens apresentam dados estatísticos similares para os diferentes canais multipercurso simulados, assim como ocorre anteriormente para largura de banda de 5MHz.

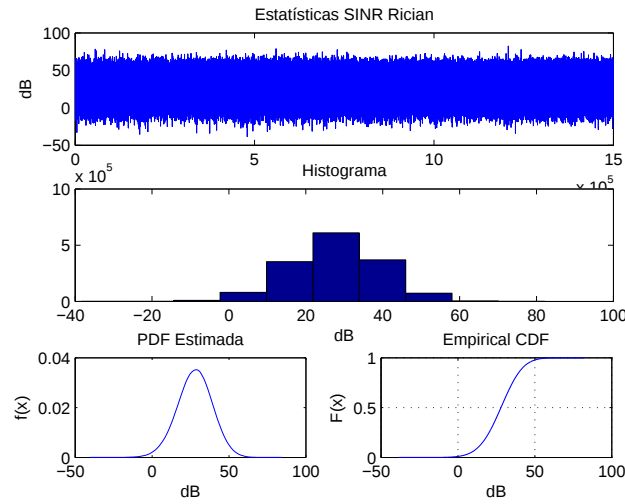
Tabela 4.12: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs

Estatística / Modelo	Rayleigh	Rician (Fator $K = 10$)	Rician (Fator $K = 100$)
Média	27.68	27.75	27.76
Desvio padrão	11.43	11.41	11.40
Variância	130.77	130.17	130.03
Valor mínimo	-40.49	-38.48	-38.51
Valor máximo	79.23	82.17	82.34
Índice de dispersão	4.72	4.69	4.68
Variância normalizada	0.17	0.16	0.16
Momento de 3ª ordem (simetria)	-187.18	-192.5	-194.67
Momento de 4ª ordem (concentração)	53467.8	52910.4	52924.5

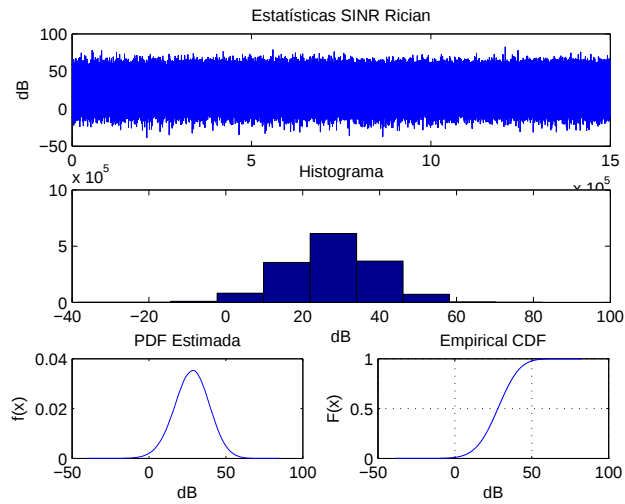
O mesmo ocorre em relação ao histograma, a PDF e a CDF das três modelagens, conforme pode ser verificado na Figura 4.17. As três modelagens apresentam comportamento similar.



(a)



(b)



(c)

Figura 4.17: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c) e 1000 TTIs

O algoritmo proposto garante taxa de atendimento ao critério de retardo máximo próximo de 100%, conforme pode ser verificado na Tabela 4.13. O algoritmo QoS garantido apresenta a menor taxa de atendimento para todos os cenários simulados. Neste cenário com largura de banda de 10MHz, os algoritmos PSO e QoS garantido apresentam melhoria em relação à taxa de atendimento, em virtude do aumento da largura de banda e consequente aumento na capacidade de transmissão.

O algoritmo PSO apresenta em geral maiores valores para vazão total e vazão média por usuário, conforme pode ser visto nas Figuras 4.18 e 4.19, fato já verificado anteriormente. Porém os valores de vazão tendem a diminuir para o algoritmo PSO a medida que aumenta o número de usuários, devido a função de penalidade característica deste algoritmo imposta pela taxa mínima. O algoritmo proposto e o algoritmo QoS garantido tendem a apresentar valores semelhantes, principalmente se considerados mais de 15 usuários em simulação, em virtude do número maior de blocos estimados para os usuários com melhores condições de canal no algoritmo proposto.

Quanto ao retardo médio por usuário, o algoritmo proposto apresenta no geral os menores valores, bastante inferiores aos valores apresentados pelo algoritmo PSO e pelo algoritmo QoS garantido, conforme verificado na Figura 4.20. O algoritmo proposto também garante a alocação de blocos de recursos para todos os usuários, independente do cenário simulado, visto que não apresenta retardo médio tendendo ao infinito, ou seja, tempo de espera infinito no *buffer*.

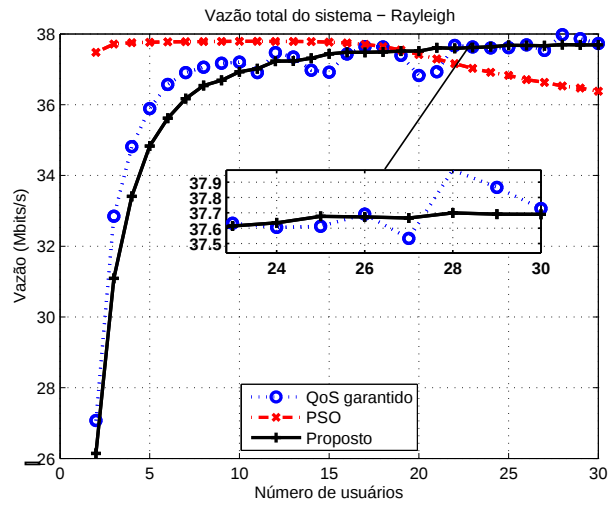
A Figura 4.21 mostra que o algoritmo proposto apresenta valores constantes e superiores à 0.8 para o índice de justiça em função da sua característica de balanceamento, enquanto os valores apresentados pelo algoritmo QoS garantido tendem a decrescer à medida que o número de usuários em simulação aumenta, em virtude de seu objetivo de atender a taxa mínima para usuários prioritários, em detrimento daqueles com menor prioridade.

Em relação a taxa de perda média do sistema, ilustrada na Figura 4.22, o algoritmo proposto apresenta melhor desempenho, com os menores valores em comparação aos algoritmos PSO e QoS garantido, devido a sua característica de balanceamento que proporciona transmissão de dados para todos os usuários no cenário simulado. O algoritmo QoS garantido apresenta no geral os maiores valores, especialmente se considerado mais de 24 usuários em simulação.

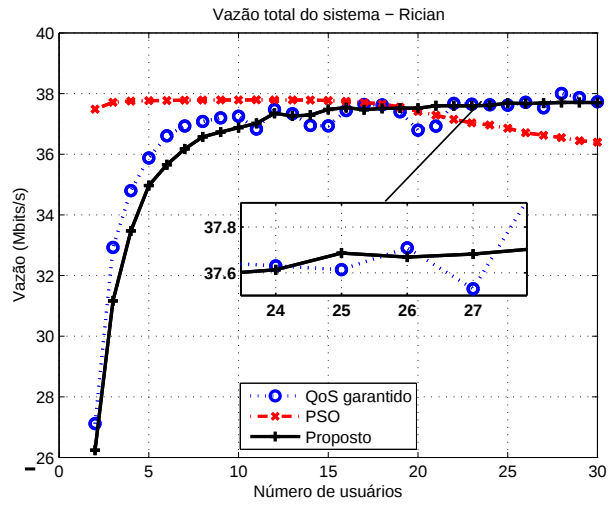
A Tabela 4.14 mostra que o algoritmo PSO apresenta maiores valores de tempo de processamento para todos os cenários simulados. O algoritmo QoS garantido e o algoritmo proposto apresentam desempenho similar em termos de tempo de processamento, com valores considerados baixos.

Tabela 4.13: *Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs*

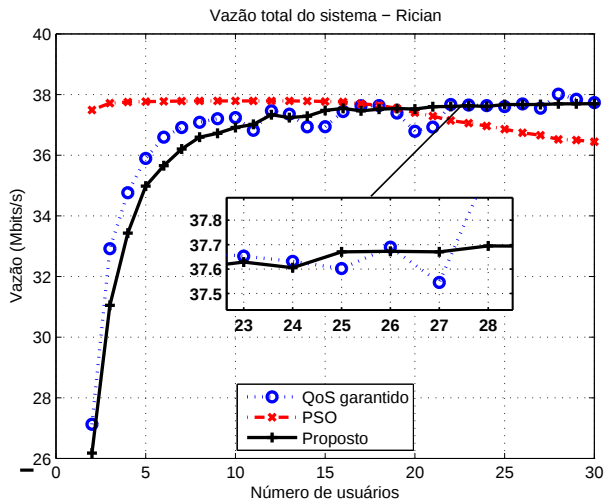
Algoritmo	Modelo multipercurso	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	Rayleigh	42688	90.80
	Rician (Fator $K = 10$)	42712	90.79
	Rician (Fator $K = 100$)	42730	90.79
PSO	Rayleigh	9850	97.87
	Rician (Fator $K = 10$)	9741	97.90
	Rician (Fator $K = 100$)	9733	97.90
Proposto	Rayleigh	1786	99.61
	Rician (Fator $K = 10$)	1739	99.62
	Rician (Fator $K = 100$)	1823	99.60



(a)

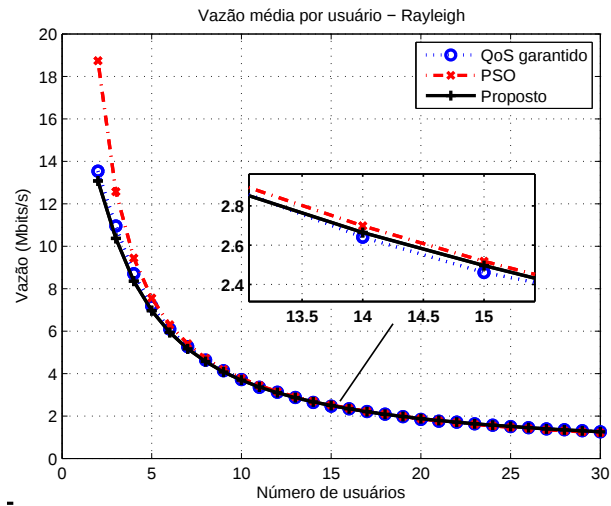


(b)

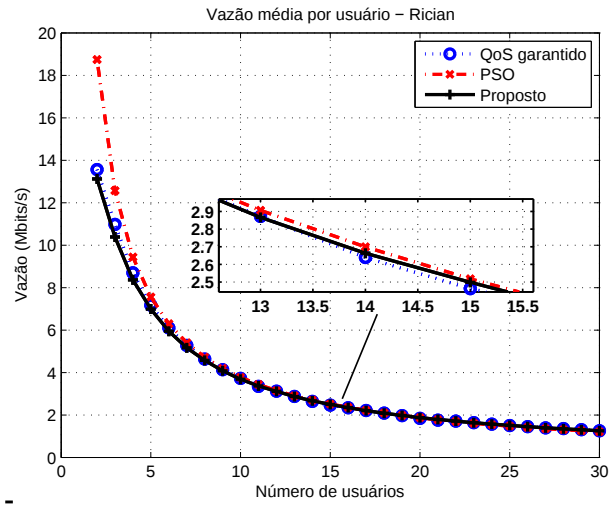


(c)

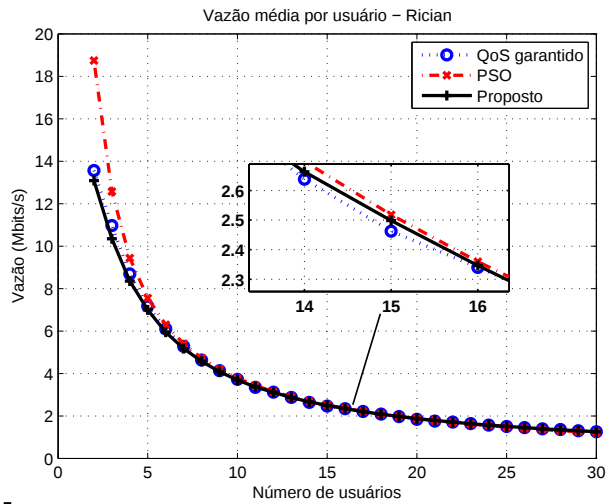
Figura 4.18: Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)

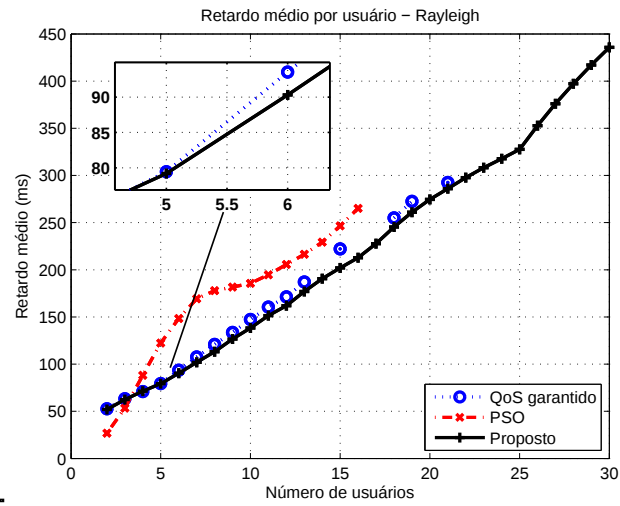


(b)

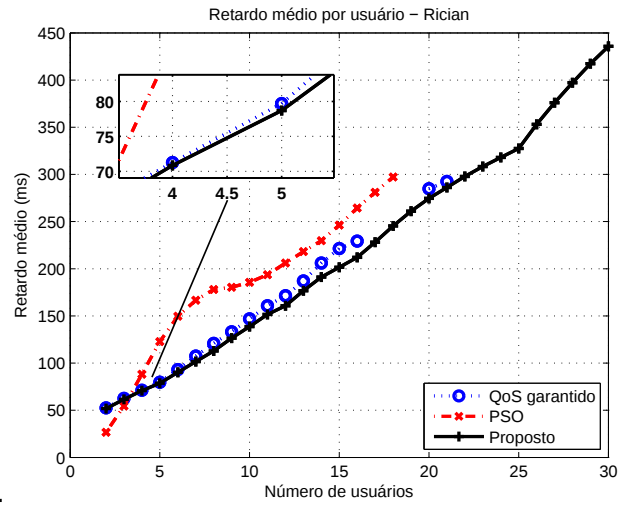


(c)

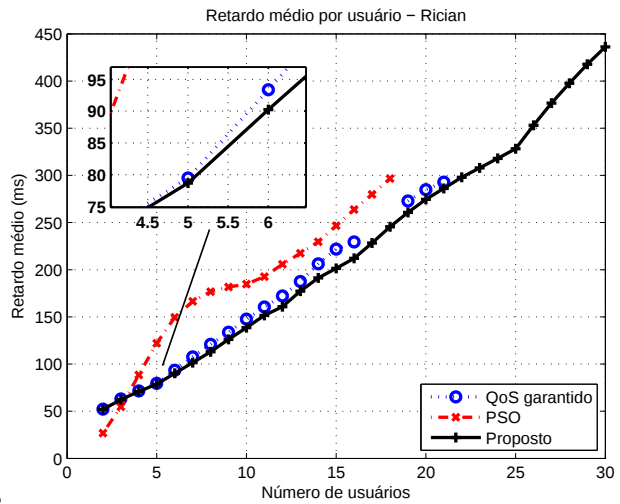
Figura 4.19: Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)

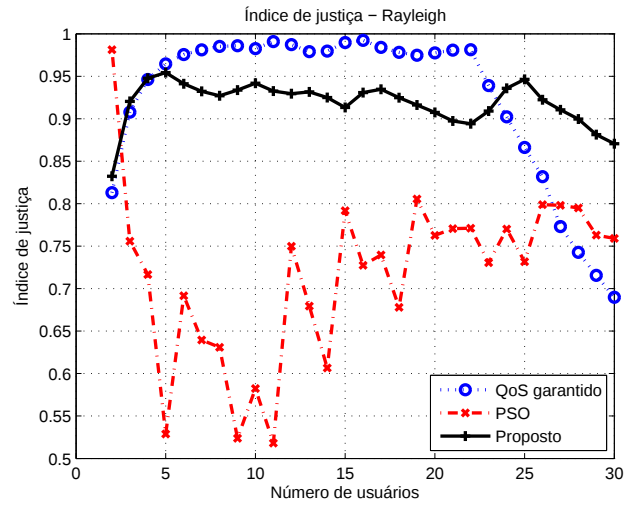


(b)

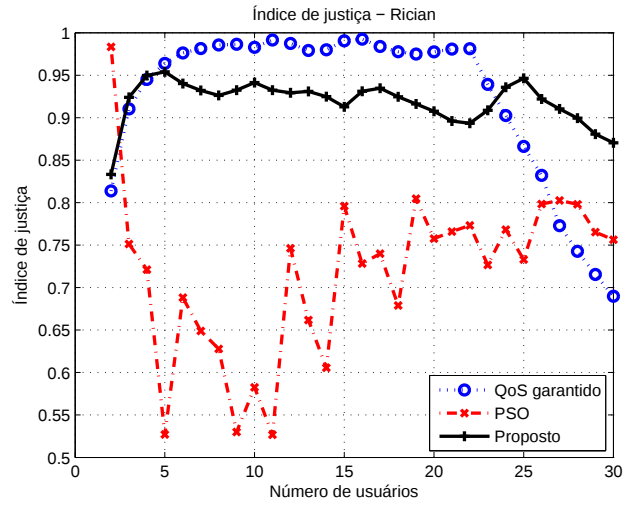


(c)

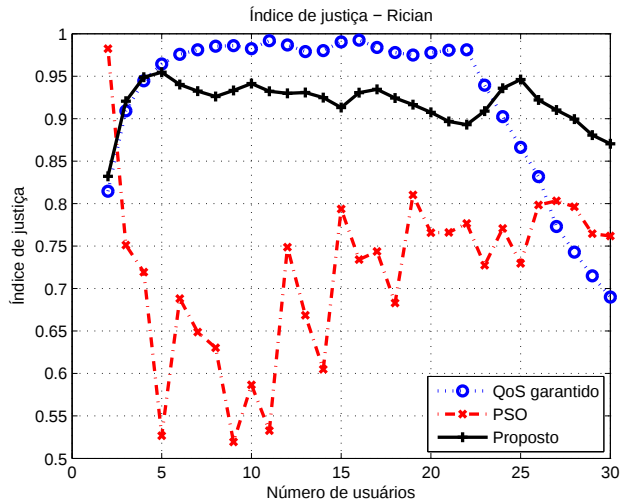
Figura 4.20: Retardo médio entre usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs



(a)



(b)



(c)

Figura 4.21: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multi-percurso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs

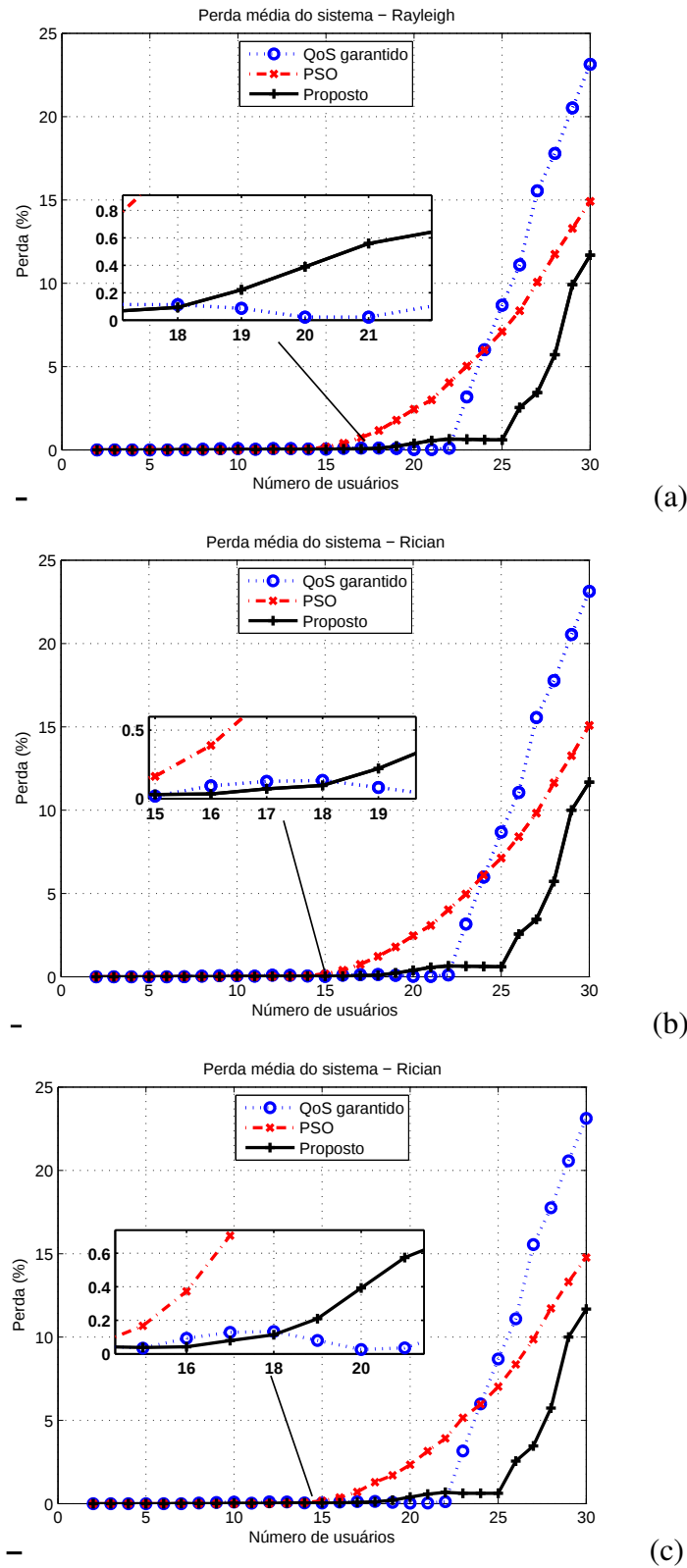


Figura 4.22: Perda média em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercorso Rayleigh (a) e Rician (com fator $K = 10$ (b) e 100 (c)) e 1000 TTIs

Tabela 4.14: *Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 0.7km, largura de banda de 10MHz, canais multipercurso Rayleigh e Rician (com fator $K = 10$ e 100) e 1000 TTIs*

Algoritmo	Modelo	Tempo de processamento (ms)				
		6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS	Rayleigh	0.410	0.582	0.738	0.855	1.025
	Rician (Fator $K = 10$)	0.406	0.586	0.742	0.859	1.032
	Rician (Fator $K = 100$)	0.414	0.595	0.743	0.860	1.038
PSO	Rayleigh	78.554	82.429	86.342	89.873	93.209
	Rician (Fator $K = 10$)	76.940	80.927	84.912	88.512	92.152
	Rician (Fator $K = 100$)	78.409	82.448	86.021	89.587	93.232
Proposto	Rayleigh	0.448	0.637	0.739	0.961	1.018
	Rician (Fator $K = 10$)	0.426	0.626	0.726	0.945	1.013
	Rician (Fator $K = 100$)	0.440	0.644	0.741	0.959	1.032

4.3.5 Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 24 Horas de Simulação

Nesta subseção considerou-se uma distância entre usuário e estação base de 0.7km, largura de banda de 10MHz, taxa mínima fixa e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 30 usuários em 864×10^5 TTIs, correspondente a um período de 24 horas.

Neste cenário o algoritmo proposto apresenta taxa de atendimento ao critério próximo de 100%, conforme verificado na Tabela 4.15. Os valores de taxa de atendimento são similares aos valores apresentados na Tabela 4.13 da subseção anterior.

O algoritmo proposto e o algoritmo QoS garantido apresentam desempenho similar em termos de vazão total e vazão média, conforme verificado nas Figuras 4.23 e 4.24, e mantém valores semelhantes aos valores apresentados na subseção anterior.

O mesmo ocorre em relação ao retardo médio entre usuários, índice de justiça e taxa de perda, mostrados nas Figuras 4.25, 4.26 e 4.27. O desempenho em termos destes parâmetros se mantém inalterado se comparado com o cenário de 1000 TTIs.

O tempo de processamento apresentado na Tabela 4.16 também é semelhante ao tempo de processamento com 1000 TTIs, mostrado na Tabela 4.14.

Tabela 4.15: *Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação*

Algoritmo	Taxa (%)
QoS garantido	90.78
PSO	97.86
Proposto	99.59

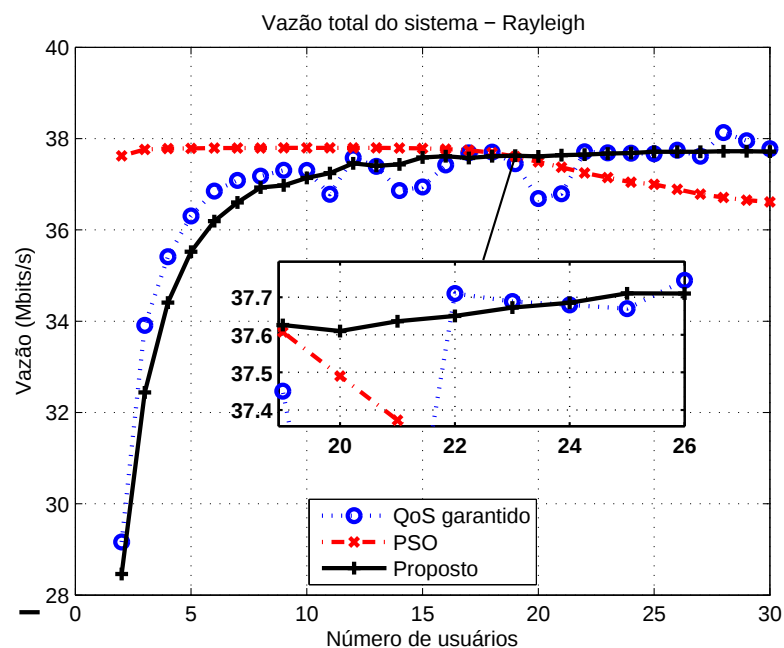


Figura 4.23: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercorso Rayleigh e 24 horas de simulação

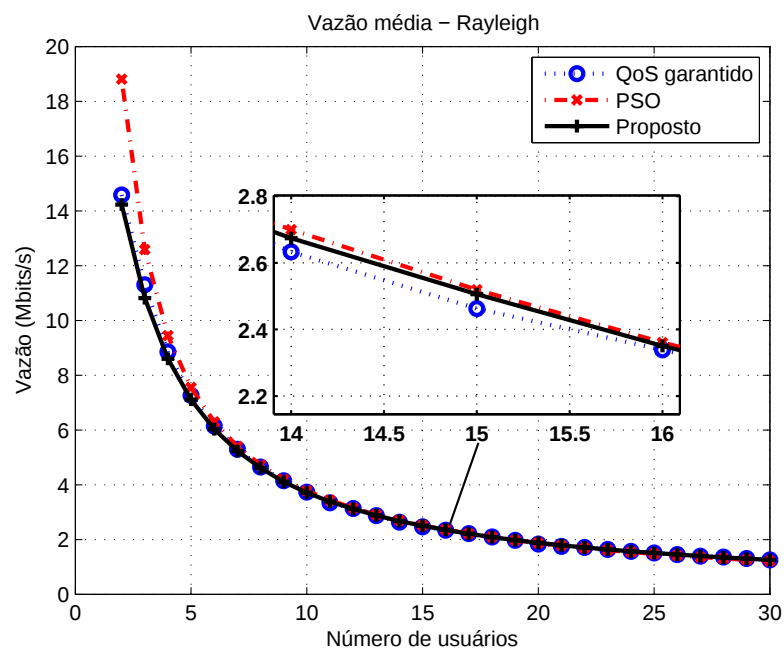


Figura 4.24: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercorso Rayleigh e 24 horas de simulação

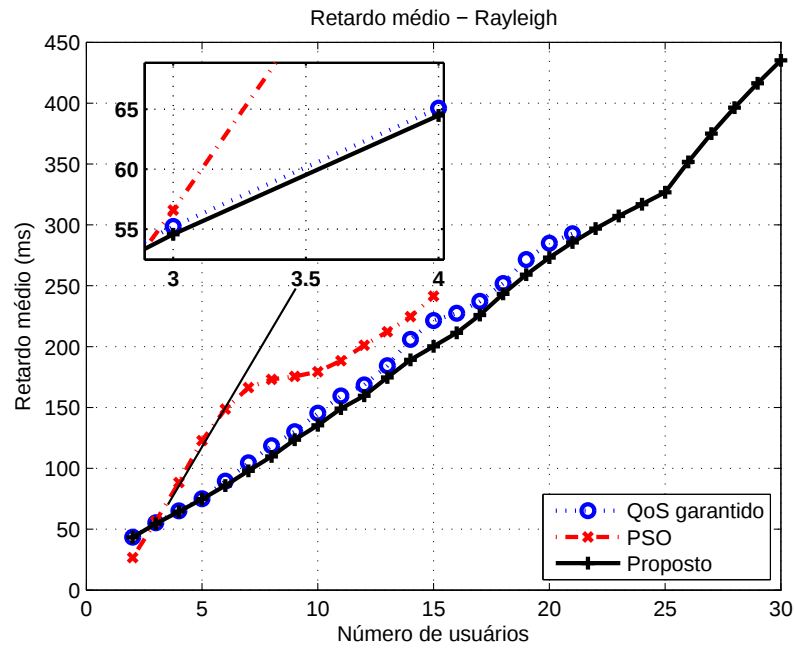


Figura 4.25: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

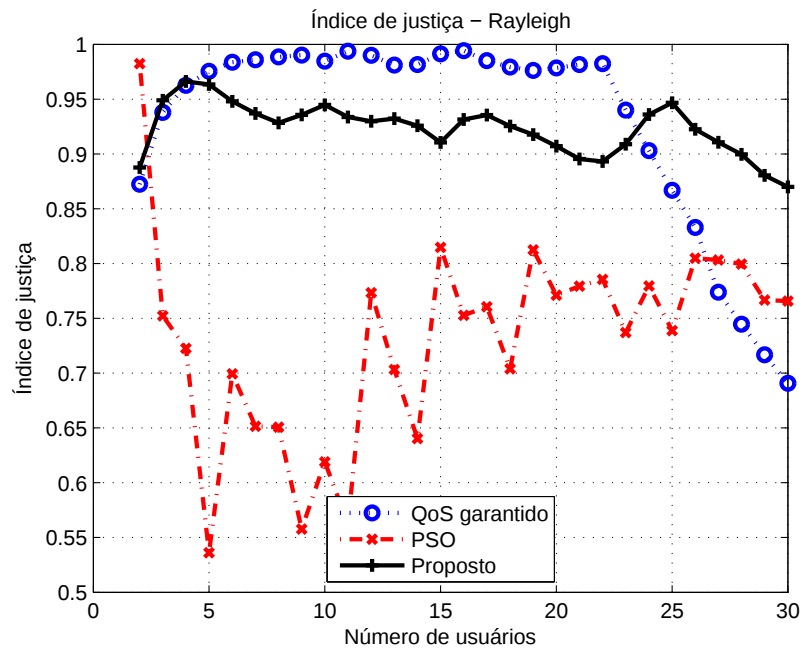


Figura 4.26: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

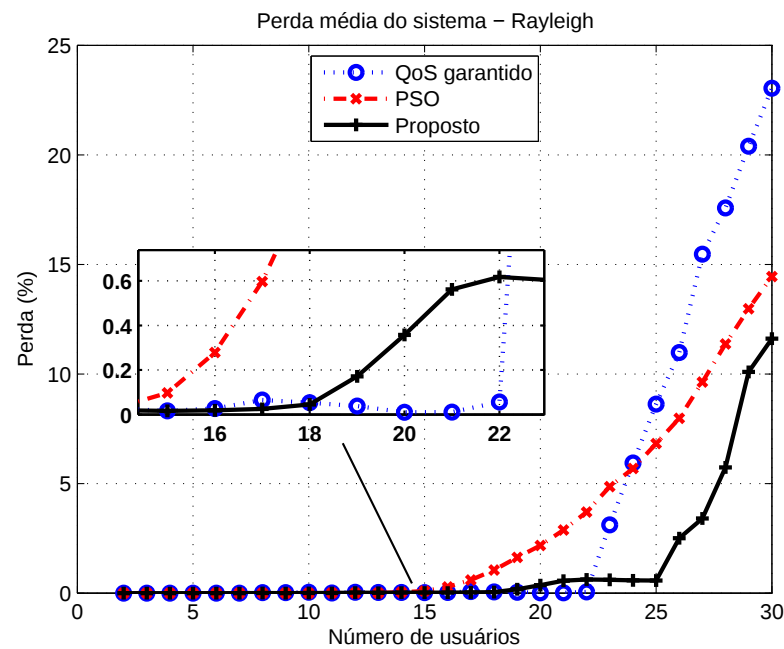


Figura 4.27: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

Tabela 4.16: Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

Algoritmo	Tempo de processamento (ms)				
	6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS garantido	0.421	0.602	0.761	0.882	1.058
PSO	79.539	83.708	87.548	91.162	94.576
Proposto	0.437	0.652	0.739	0.962	1.028

4.3.6 Distância de 0.7km, Largura de Banda de 10MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs

Nesta subseção considerou-se uma distância entre usuário e estação base de 0.7km, largura de banda de 10MHz, taxa mínima calculada adaptativamente conforme explicado na Seção 3.5 e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 30 usuários em 1000 TTIs.

O algoritmo proposto apresenta taxa de atendimento ao critério de retardo máximo igual a 99.10%, maior que a taxa apresentada pelos algoritmos QoS garantido e PSO, conforme pode ser verificado na Tabela 4.17.

Para mais de 14 usuários em simulação, o algoritmo proposto apresenta em geral vazão total maior que a apresentada pelo algoritmo QoS garantido, conforme ilustrado na Figura 4.28. O algoritmo PSO apresenta em geral maior vazão total, sem tendência de decrescer o valor, pois neste cenário a função de penalidade imposta pela taxa mínima tende a ser menor do que no cenário com taxa mínima fixa. O fato se repete em relação à vazão média, apresentada na Figura 4.29.

A Figura 4.30 mostra que o algoritmo proposto apresenta em geral menor retardo médio entre usuários, enquanto o algoritmo PSO apresenta os maiores valores de retardo. O algoritmo proposto é o único que garante transmissão para todos os usuários neste cenário, ou seja, não apresenta retardo tendendo ao infinito.

Em relação ao índice de justiça, ilustrado na Figura 4.31, o algoritmo QoS garantido apresenta melhor desempenho, e o algoritmo PSO apresenta os menores valores.

Quanto a taxa de perda média do sistema, ilustrada na Figura 4.32, o algoritmo PSO apresenta os maiores valores. O algoritmo proposto apresenta valores menores do que 1%, similares aos valores apresentados pelo algoritmo QoS garantido. A característica de balanceamento do algoritmo proposto tende a promover melhoria na taxa de perda. Em relação ao algoritmo QoS garantido, os valores menores de taxa mínima proporcionados pelo cálculo adaptativo também tendem a promover melhoria neste quesito, uma vez que o critério é mais fácil de ser atingido.

A Tabela 4.18 mostra que o algoritmo proposto e o algoritmo QoS garantido apresentam valores similares de tempo de processamento, enquanto o algoritmo PSO apresenta pior desempenho.

Tabela 4.17: *Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs*

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	47909	89.67
PSO	89609	80.69
Proposto	4153	99.10

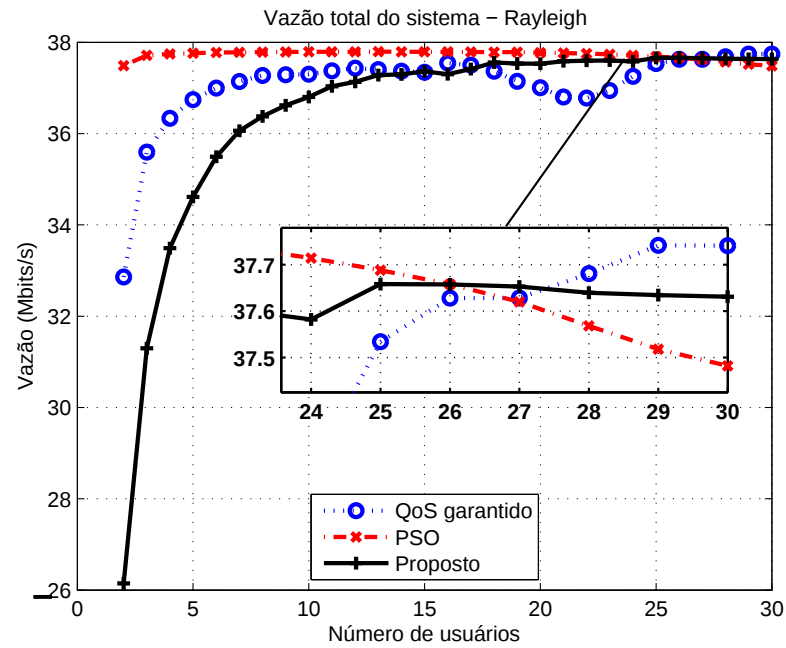


Figura 4.28: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

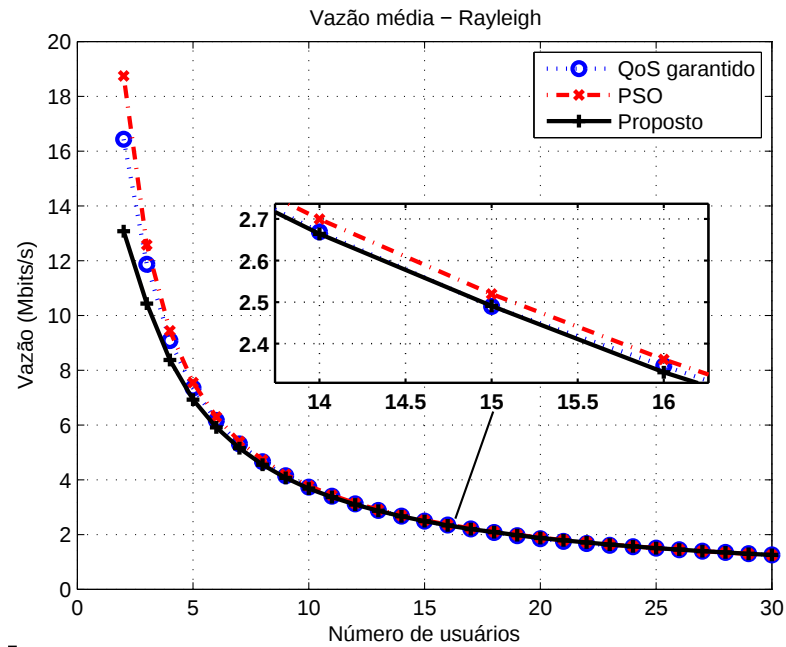


Figura 4.29: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

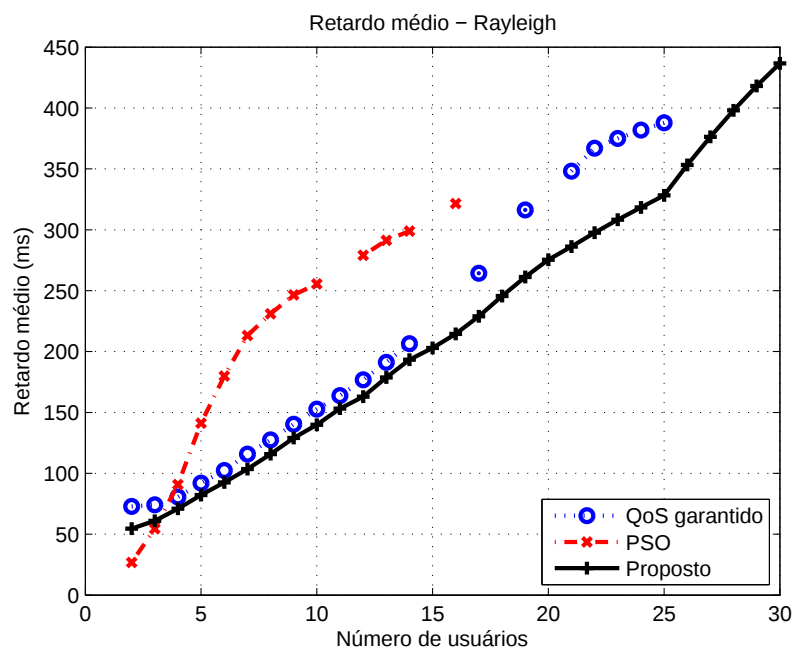


Figura 4.30: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

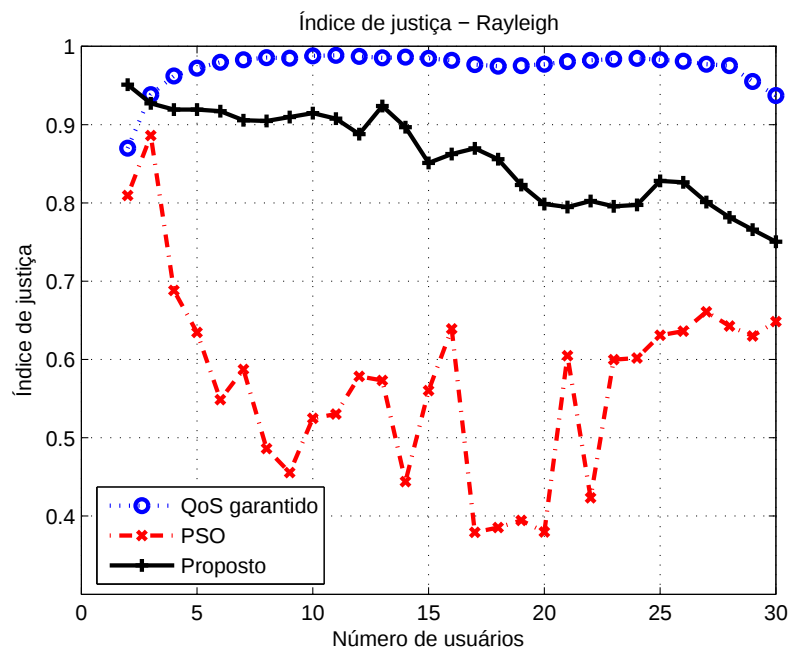


Figura 4.31: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

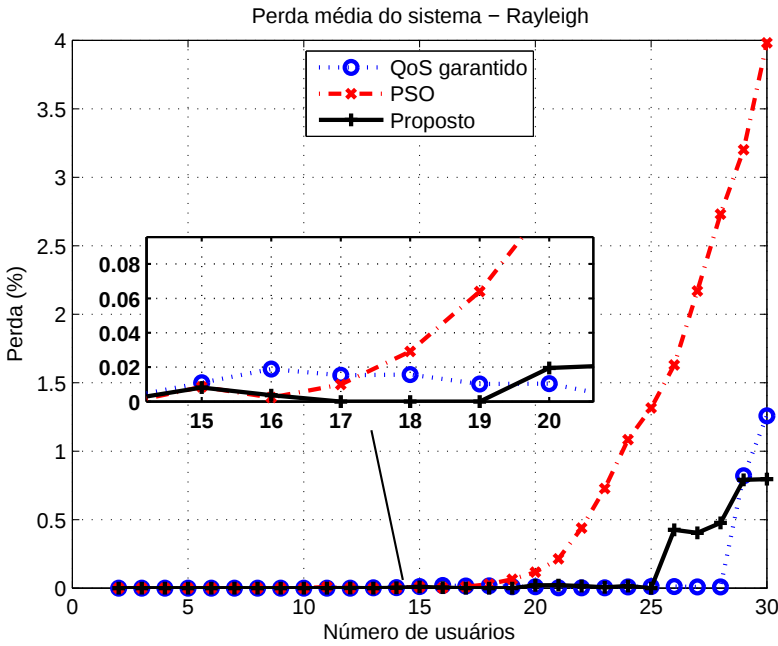


Figura 4.32: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 4.18: Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Tempo de processamento (ms)				
	6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS garantido	0.429	0.596	0.748	0.881	1.058
PSO	81.420	84.954	89.143	92.635	96.263
Proposto	0.437	0.645	0.746	0.962	1.028

Alocação de Recursos em Redes LTE com Estimação de Retardo Utilizando Curva de Serviço e Processo Envelope

Neste capítulo é proposto um algoritmo de alocação de blocos de recurso para sistemas de comunicação LTE (*Long Term Evolution*) no qual são levadas em conta as restrições impostas pelo esquema MCS (*Modulation and Coding Scheme*) para transmissão *downlink*. O sistema proposto considera a qualidade de transmissão do canal, o critério de retardo máximo e a estimativa de retardo para se decidir sobre a alocação de recursos de rádio disponíveis, buscando reduzir o retardo médio do sistema. Para a estimativa de retardo são utilizados conceitos de cálculo de rede determinístico, curva de serviço [13] e processo envelope MFBAP (*Multifractal Bounded Arrival Process*) [44]. Através da estimativa de limitante de retardo, pretende-se prover subsídios para tomada de decisão sobre o controle de admissão de usuários a fim de garantir valor de retardo máximo estabelecido.

São realizadas comparações com outros algoritmos de alocação de recursos através de parâmetros de QoS (*Quality of Service*) como retardo médio, vazão, tempo de processamento, taxa de perda e índice de justiça (*fairness*), comprovando a eficiência do algoritmo proposto [20].

5.1 Cálculo de Rede Determinístico

O Cálculo de Rede Determinístico pode ser utilizado para estimar recursos a fim de prover QoS em redes e tem fornecido ferramentas poderosas para estimação do *backlog* e retardo em uma rede com garantia de serviço para fluxos de tráfego individuais. Usando a noção de processo envelope, curvas de chegada e curvas de serviço, vários trabalhos tem demonstrado que os limitantes de *backlog* e retardo podem ser concisamente expressos pela álgebra Min-Plus [9].

O Cálculo de Rede também pode ser visto como a teoria de sistemas que se aplica às redes de computadores, mas a principal diferença é considerar-se outra álgebra, onde as operações são alteradas da seguinte forma: adição torna-se o cálculo do mínimo, e a multiplicação torna-se adição [9].

Seja S um subconjunto não vazio de $\mathbb{R}$. S é o limite inferior se houver um número M tal que $s \geq M$ para todo $s \in S$. Adotou-se chamar de ínfimo de S , e denotada por $\inf S$. Por exemplo, os intervalos fechados e abertos $[a, b]$ e (a, b) tem o mesmo ínfimo, que é a . Agora, se S contém um elemento que é menor do que todos os seus outros elementos, esse elemento é chamado de mínimo de S , e é denotado por $\min S$. Note que o mínimo de um conjunto, nem sempre existe. Por exemplo, (a, b) não tem mínimo desde que $a \notin (a, b)$. Por outro lado, se o mínimo de um conjunto S existe, é idêntico ao seu ínfimo. Por exemplo, $\min[a, b] = \inf[a, b] = a$. É facilmente mostrado que todo subconjunto não vazio finito de $\mathbb{R}$ tem um mínimo [9].

Seja um fluxo de entrada $x(t)$ e um fluxo de saída $y(t)$ que se conforma a um conjunto de taxas de acordo com um processo envelope de tráfego A (série de tráfego), à custa de, possivelmente, atrasar bits no *buffer*. Essas funções são não-decrescentes com o tempo t , seja contínuo ou discreto. A convolução Min-Plus entre A e x é dada por [9]:

$$y(t) = (A \otimes x)(t) = \inf_{s \in \mathbb{R} | 0 \leq s \leq t} \{A(t-s) + x(s)\}. \quad (5-1)$$

A convolução na teoria de sistemas tradicionais é tanto comutativa quanto associativa. Por exemplo, a resposta ao impulso do circuito da Figura 5.1 (a) é a convolução da resposta ao impulso de cada uma das células elementares [9]:

$$h(t) = (h_1 \otimes h_2)(t) = \int_0^t h_1(t-s)h_2(s)ds. \quad (5-2)$$

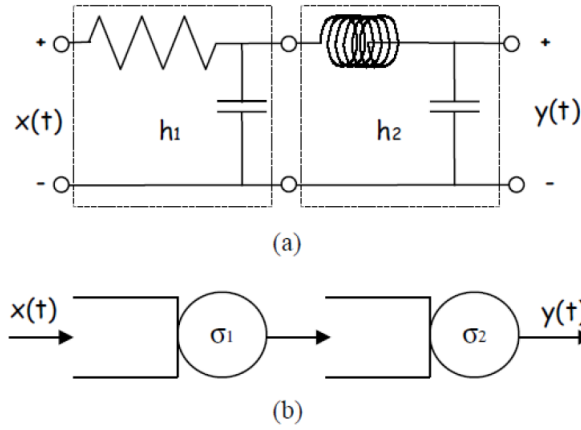


Figura 5.1: A resposta ao impulso da concatenação de dois circuitos lineares é (a) a convolução da resposta ao impulso individual, a curva de formação da concatenação de dois reguladores é (b) a convolução das curvas individuais modeladas [9]

A mesma propriedade é utilizada para o regulador de tráfego. A saída do segundo regulador na Figura 5.1 (b) é na verdade igual a $y(t) = (A \otimes x)(t)$, onde [9]:

$$A(t) = (A_1 \otimes A_2)(t) = \inf_{s \in \mathbb{R} | 0 \leq s \leq t} \{A_1(t-s) + A_2(s)\}. \quad (5-3)$$

O tráfego de rede pode ser descrito utilizando-se funções de entrada $R(t)$ e saída $R^*(t)$, obtendo assim o retardo (*delay*) no instante t , ou seja, o retardo experimentado por um bit chegando no tempo t se todos os bits recebidos antes dele forem servidos antes também. O retardo é dado por [9]:

$$d(t) = \inf \{\tau \geq 0 : R(t) \leq R^*(t + \tau)\}. \quad (5-4)$$

5.1.1 Processo envelope

O processo envelope para o tráfego de chegada de pacotes é um limitante superior para o processo real de tráfego de pacotes acumulados. Para um processo envelope determinístico, a função limitante $\hat{A}(t)$ corresponde ao valor máximo de um fluxo $A(t)$ no intervalo de tempo $[s, s+t]$, e é definida pela equação [9]:

$$\hat{A}(t) = \sup_{s \geq 0} A[s, s+t]. \quad (5-5)$$

Balde Furado Fractal (FLB, *Fractal Leaky Bucket*)

O FLB é um mecanismo de policiamento introduzido por [32], baseado no modelo fBm (*fractional Brownian motion*) de processo de tráfego de pacotes. Foi demons-

trado que o fBm descreve com precisão fluxos de tráfego, dado sua média ($\bar{a}$), desvio padrão (σ) e parâmetro de Hurst (H) [32].

A quantidade máxima de pacotes policiados, ou seja, o seu processo envelope, é dado pela equação [32]:

$$\hat{A}_{FLB}(t) = \bar{a}t + k\sigma t^H + B, \quad (5-6)$$

onde H é o parâmetro de Hurst, t é o instante de tempo, $\bar{a}$ e σ são respectivamente, a média e o desvio padrão do tráfego de entrada, k é a constante relacionada à probabilidade de violação (para $\varepsilon = 10^{-6}$) do processo envelope e B é o tamanho do *buffer*.

Processo de Chegada com Limitante Multifractal (MFBAP - *Multifractal Bounded Arrival Process*)

O MFBAP é uma alternativa determinística de se obter o processo envelope que limita o volume do tráfego em um dado intervalo de tempo, calculado da seguinte forma [44]:

$$\hat{A}_{MFBAP}(t) = \bar{a}t + k\sigma t^{H(t)} + B, \quad (5-7)$$

onde $H(t)$ é o expoente de Hölder, t é o instante de tempo, $\bar{a}$ e σ são respectivamente, a média e o desvio padrão do tráfego de entrada, k é a constante relacionada à probabilidade de violação (para $\varepsilon = 10^{-6}$) do processo envelope e B é o tamanho do *buffer*.

Análise Comparativa

As Figuras 5.2, 5.3, 5.4, 5.5 e 5.6 apresentam o processo envelope MFBAP, calculado conforme equação 5-7, o processo envelope FLB, calculado conforme equação 5-6, e o processo envelope real, isto é, o tráfego acumulado, calculado conforme equação 5-5, para as séries reais de tráfego TCP/IP [48], apresentadas no Apêndice A. O expoente de Hölder $H(t)$ foi calculado utilizando a ferramenta FracLab 2.0 [22], através do método baseado em oscilações. Pode ser observado pelos gráficos que o processo envelope FLB tende a superestimar os valores de envelope, enquanto o processo envelope MFBAP apresenta no geral valores similares aos apresentados pelo processo envelope real, portanto com maior precisão. O fato se justifica pois o expoente de Hölder é calculado de forma determinística e variante no tempo, enquanto o parâmetro de Hurst fornece um valor fixo que pode não descrever com precisão as características de diferentes fluxos de tráfego.

A escolha pelo processo envelope MFBAP para estimativa do limitante de retardo se justifica pois o MFBAP é uma heurística que se baseia na modelagem multifractal que apresenta resultados mais próximos ao envelope real do que o FLB, monofractal,

apesar de ambos levarem em consideração o tráfego em rajadas com propriedades fractais tais como impulsividade, auto-similaridade e dependência de longa duração em diversas escalas de agregação temporal, na faixa de milissegundos a minutos [31].

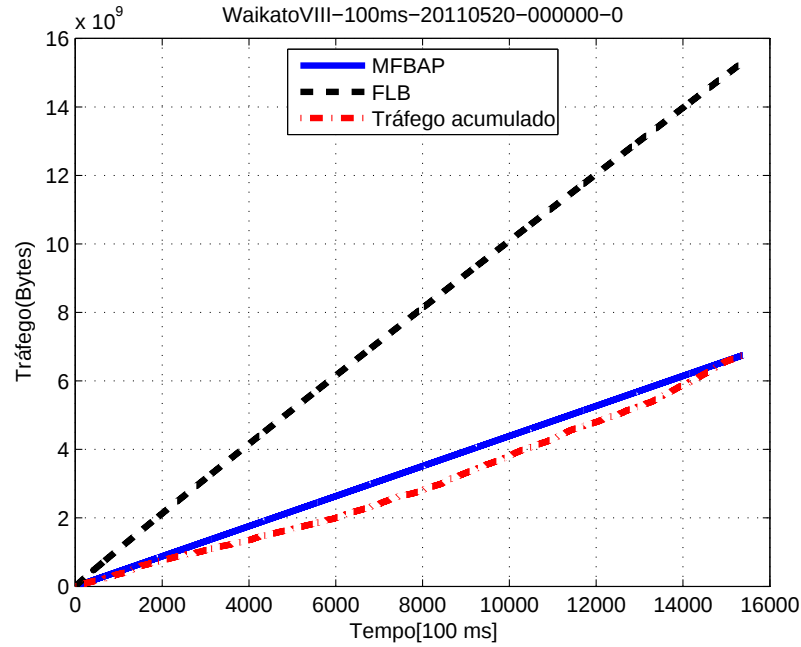


Figura 5.2: Processo envelope MFBAP, FLB e real para representação dos usuários 1, 5, 10, 15, 20, 25 e 30

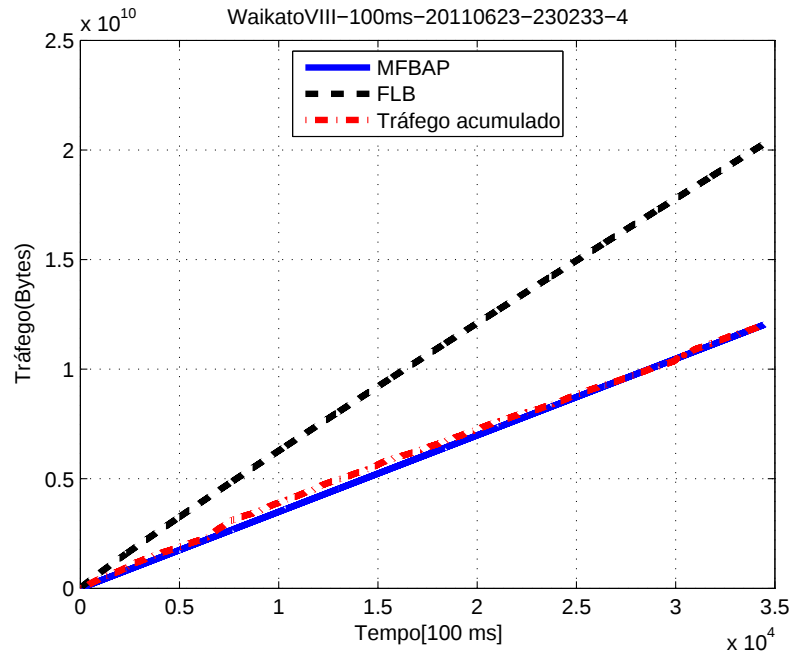


Figura 5.3: Processo envelope MFBAP, FLB e real para representação dos usuários 2, 6, 11, 16, 21 e 26

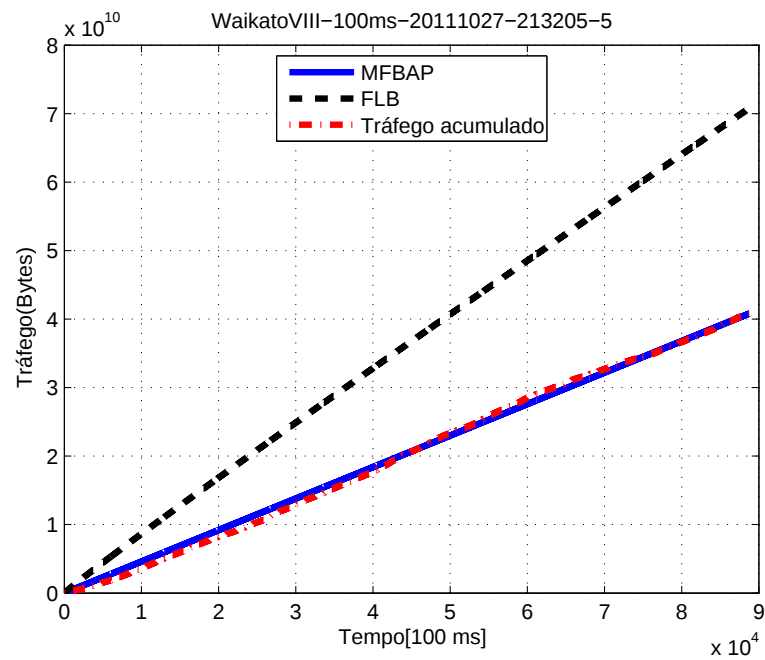


Figura 5.4: Processo envelope MFBAP, FLB e real para representação dos usuários 3, 7, 12, 17, 22 e 27

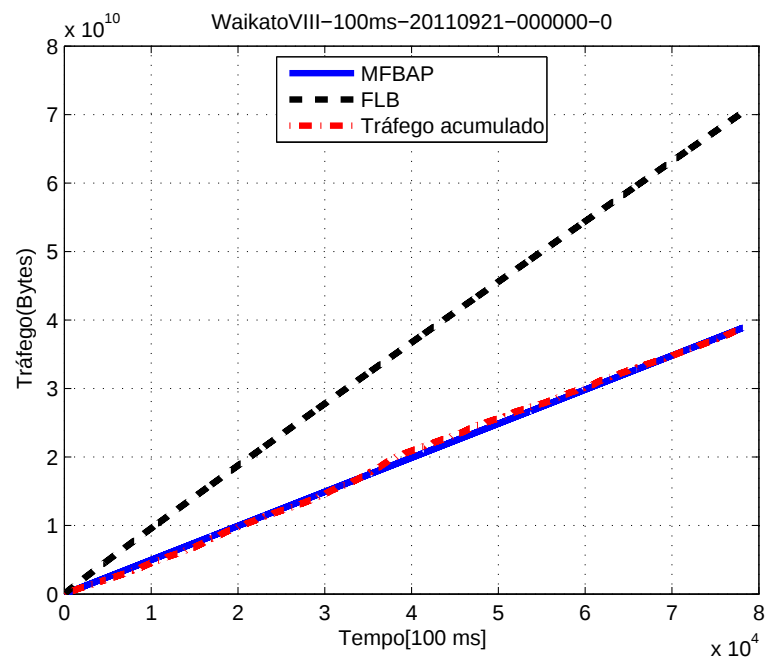


Figura 5.5: Processo envelope MFBAP, FLB e real para representação dos usuários 4, 8, 13, 18, 23 e 28

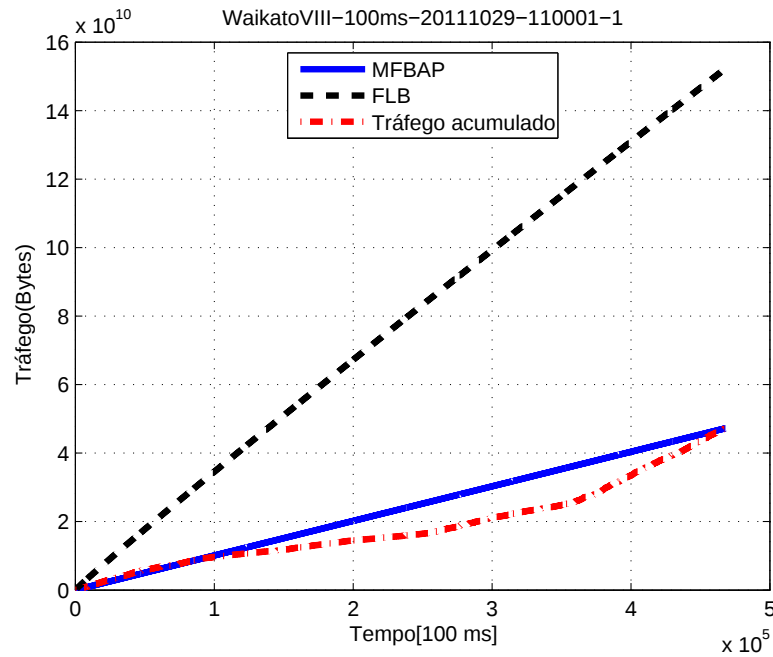


Figura 5.6: Processo envelope MFBAP, FLB e real para representação dos usuários 5, 9, 14, 19, 24 e 29

5.1.2 Curva de Serviço

Considerando que as chegadas de tráfego em um único nó de uma rede, no intervalo de tempo $[0, t)$, são dadas em termos da função $A(t)$. As saídas de um fluxo a partir do nó, no intervalo de tempo $[0, t)$, são denotadas por $D(t)$ com $D(t) \leq A(t)$.

Seja a convolução denotada por $*$ e a deconvolução por $\otimes$, uma curva mínima de serviço para um fluxo é uma função S que especifica um limite inferior para o serviço prestado ao fluxo de tal forma que, para todo $t \geq 0$ [9] [13],

$$D(t) \geq A * S(t). \quad (5-8)$$

Uma curva de serviço máxima para um fluxo é uma função $\bar{S}$ que especifica um limite superior sobre o serviço prestado a um fluxo tal que, para todo $t \geq 0$ [9] [13],

$$D(t) \leq A * \bar{S}(t). \quad (5-9)$$

A curva de serviço desempenha um importante papel no Cálculo de Rede, uma vez que ela provê garantias de serviço.

Assumindo que um sistema OFDM-TDMA (*Orthogonal Frequency-Division Multiple Access/Time Division Multiple Access*) com escalonamento Round-Robin possua 4 usuários, a curva de serviço do usuário 1 (o primeiro a ser atendido, em $t = 0$) pode ser visualizada na Figura 5.7 [13].

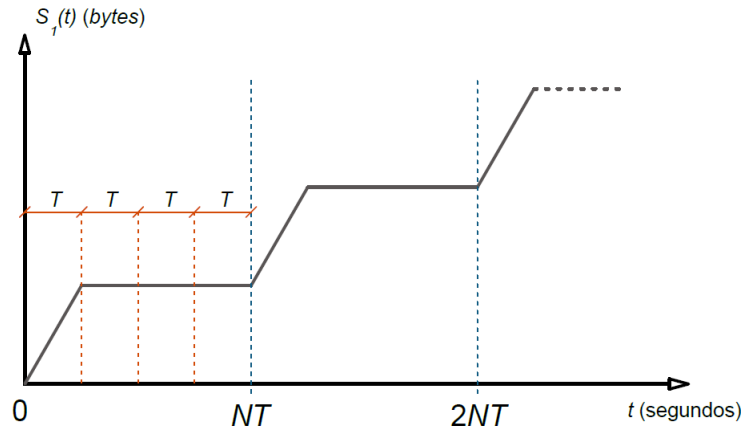


Figura 5.7: Curva de serviço do usuário 1 [13]

A ilustração mostra dois ciclos de escalonamento. Observa-se que nos primeiros T segundos de cada ciclo o usuário 1 é atendido, enquanto nos $3T$ segundos restantes ele aguarda [13].

Assume-se que o coeficiente angular da reta que vai de $t = 0$ até $t = T$ seja dado por c , a média da taxa de atendimento no servidor do sistema OFDM [13].

$$c = \sum_{r=0}^R r [R_n]_{r+1}. \quad (5-10)$$

Dessa forma, a curva de serviço do usuário 1 pode ser calculada como a soma do tráfego transmitido nos ciclos completos e do tráfego já transmitido no ciclo atual [13].

$$S_1(t) = S_{1C}(t) + S_{1A}(t). \quad (5-11)$$

O número de ciclos completos é dado por $P = \lfloor \frac{t}{NT} \rfloor$. Assim, o tráfego transmitido nos ciclos já completados é calculado por [13]:

$$S_{1C}(t) = cTP. \quad (5-12)$$

O tráfego transmitido no ciclo atual, por sua vez, é calculado por [13]:

$$S_{1A}(t) = cT \min\left(\frac{t - PNT}{T}; 1\right). \quad (5-13)$$

Logo, a curva de serviço do usuário 1 é dada por [13]:

$$S_1(t) = cTP + cT \min\left(\frac{t - PNT}{T}; 1\right). \quad (5-14)$$

Sendo assim, a curva de serviço de um sistema OFDM-TDMA com escalonamento Round-Robin generalizada para qualquer usuário servido pelo mesmo intervalo de tempo T é dado por [13]:

$$S_n(t) = cTP + cT \min \left\{ \frac{\max [t - PNT - (n-1)T; 0]}{T}; 1 \right\}. \quad (5-15)$$

Propõe-se utilizar essa curva de serviço, onde c é a média da taxa de atendimento no servidor do sistema e N é o número de intervalos de tempo T por ciclo completo P , sendo $P = \lfloor \frac{t}{NT} \rfloor$.

A Figura 5.8 apresenta diferentes curvas de serviços calculadas conforme equação 5-15 para cenários com diversos números de usuários e taxa de transmissão de dados. Nota-se que a curva de serviço apresenta valores cada vez mais elevados a medida que a taxa de transmissão aumenta ou o número de usuários diminui.

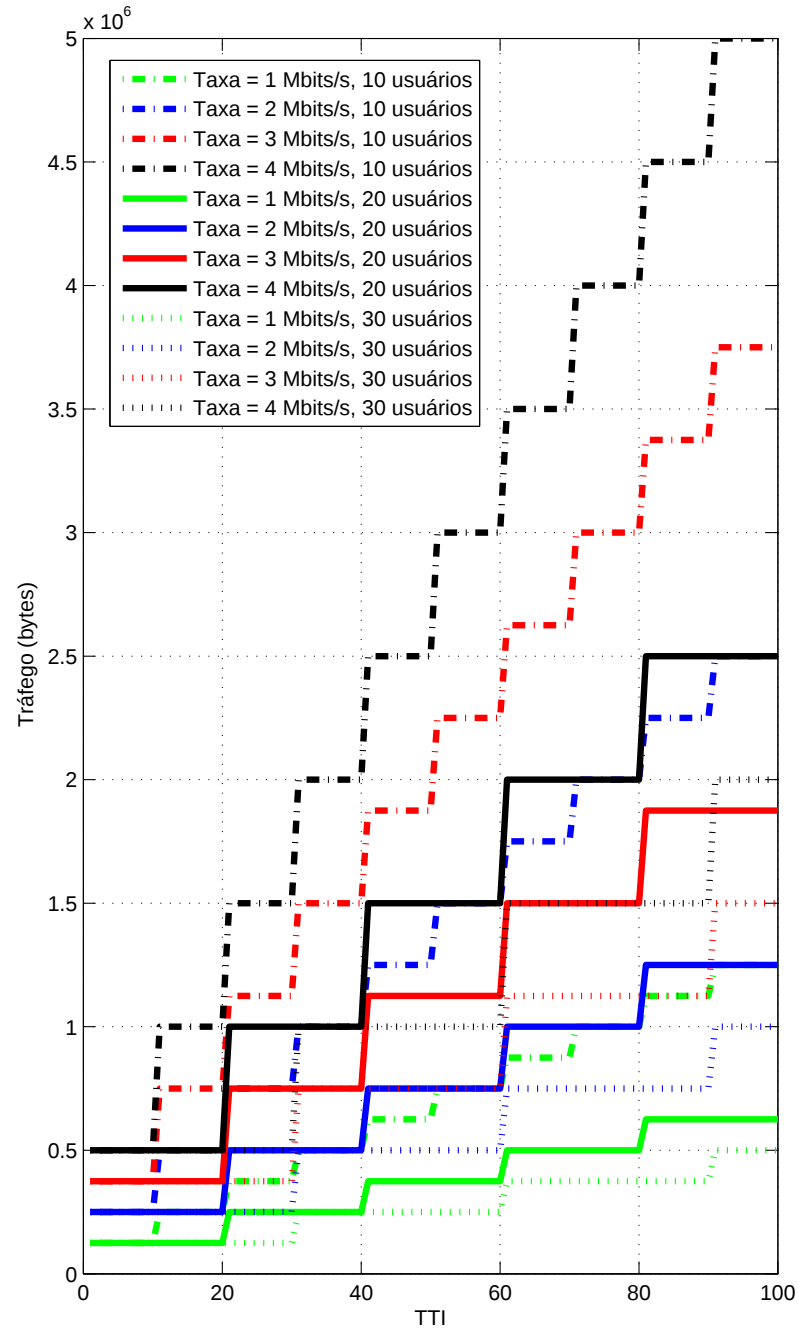


Figura 5.8: Curva de serviço por usuário calculada para diferentes cenários em 100 TTIs

5.1.3 Estimativa de Limitante de Retardo

Segundo o Cálculo de Rede Determinístico, o limitante superior de retardo estimado, denotado por $\hat{d}$, é dado por [44]:

$$\hat{d} = \inf \{d \geq 0 | \forall t \geq 0 : A^*(t-d) \leq S(t)\}. \quad (5-16)$$

Assim, propõe-se neste capítulo utilizar essa equação para estimar o retardo em redes LTE, onde $\inf$ é o valor ínfimo, A^* é o processo envelope MFBAP, calculado conforme equação 5-7, e S é a curva de serviço, calculada conforme a equação 5-15.

A curva de serviço S utilizada considera um cenário simplificado de transmissão OFDM-TDMA com escalonamento Round-Robin. Sendo assim, será avaliado se o algoritmo de alocação proposto funciona da mesma forma que o escalonador Round-Robin.

As Figuras 5.9 e 5.11 apresentam o retardo real médio entre os usuários da rede utilizando o escalonador Round-Robin e o escalonador proposto, assim como a estimativa de limitante de retardo, considerando diferentes distâncias e larguras de banda. Verifica-se que quanto maior a distância entre o usuário e a estação base, maior será o retardo e a estimativa. Em relação a estimativa de retardo, os valores se aproximam dos valores apresentados pelo algoritmo proposto, independente da distância considerada, e são inferiores aos valores apresentados utilizando escalonamento Round-Robin.

Se considerado até 6 usuários em simulação com largura de banda de 5MHz e distâncias de 1 ou 1.4km, a estimativa de retardo apresenta valores em média superiores aos valores apresentados pelo algoritmo proposto. Se considerado até 13 usuários em simulação com largura de banda de 10MHz, os valores de estimativa de retardo são em média superiores aos valores apresentados pelo algoritmo proposto. Portanto, nestes casos, a estimativa de limitante de retardo fornece valores adequados que garantem que o algoritmo proposto se comporte de forma a garantir QoS em termos de retardo.

Também pode ser notado através das Figuras 5.10 e 5.12 que o Erro Quadrático Médio (EQM) é consideravelmente baixo para os cenários simulados, tendendo a aumentar a medida que o número de usuários considerados em simulação e a distância entre a UE e a BS aumentam. Isto prova que a estimativa de retardo proposta utilizando curva de serviço e processo envelope MFBAP apresenta valores bastantes similares aos valores de retardo apresentados pelo algoritmo proposto. A média do EQM foi calculada sobre 1000 instâncias para cada um dos pontos dos gráficos que representam o número de usuários considerados em rede.

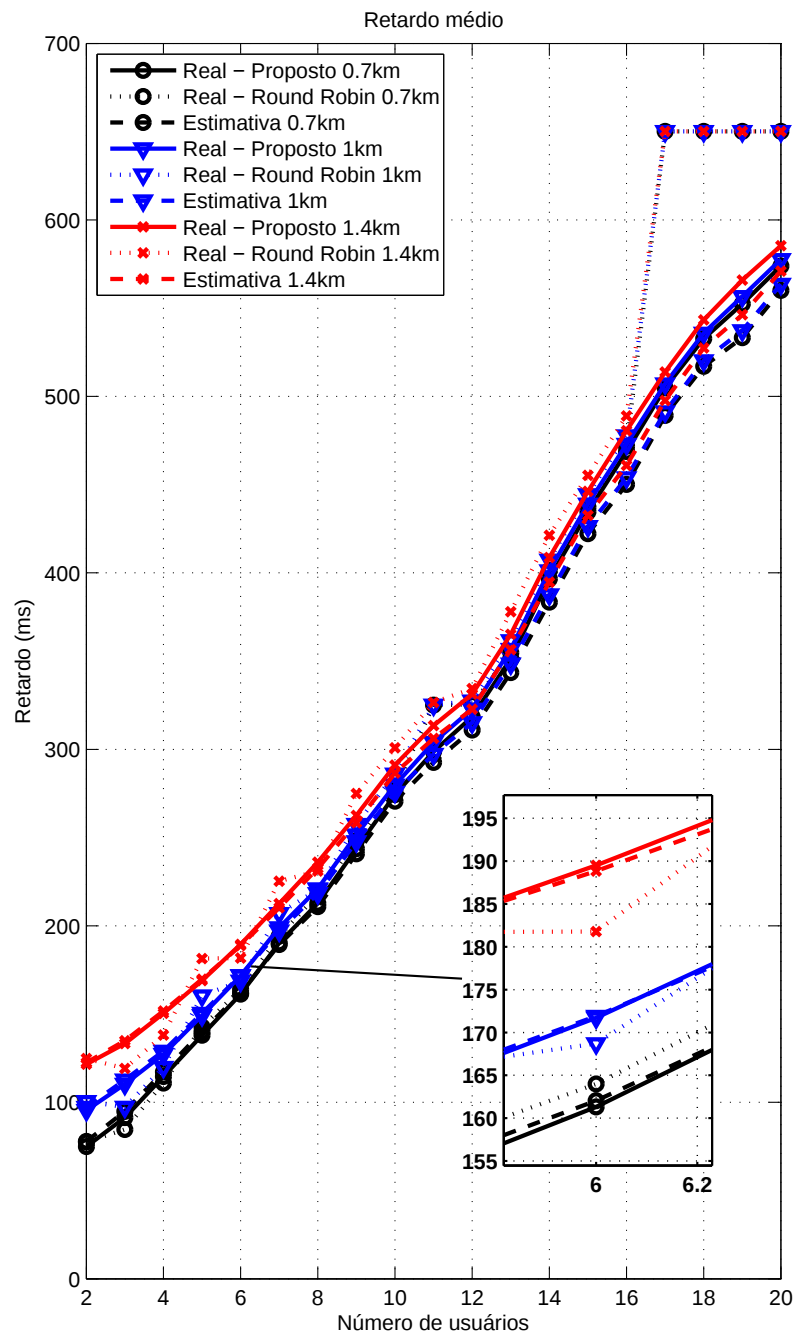


Figura 5.9: Estimativa de limitante de retardo e retardo real médio entre usuários calculado para algoritmo proposto e escalonador Round-Robin, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

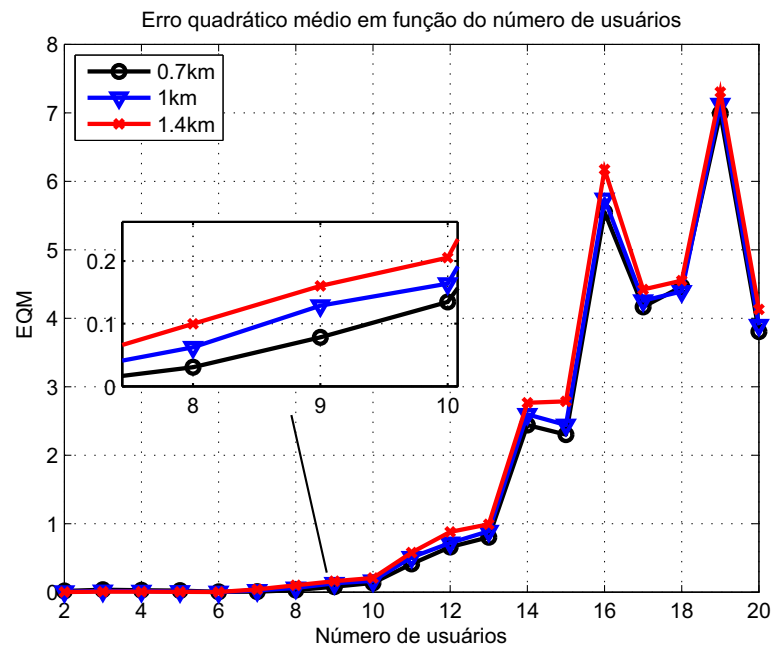


Figura 5.10: Erro Quadrático Médio (EQM) na estimativa de limitante de retardo em função do número de usuários, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

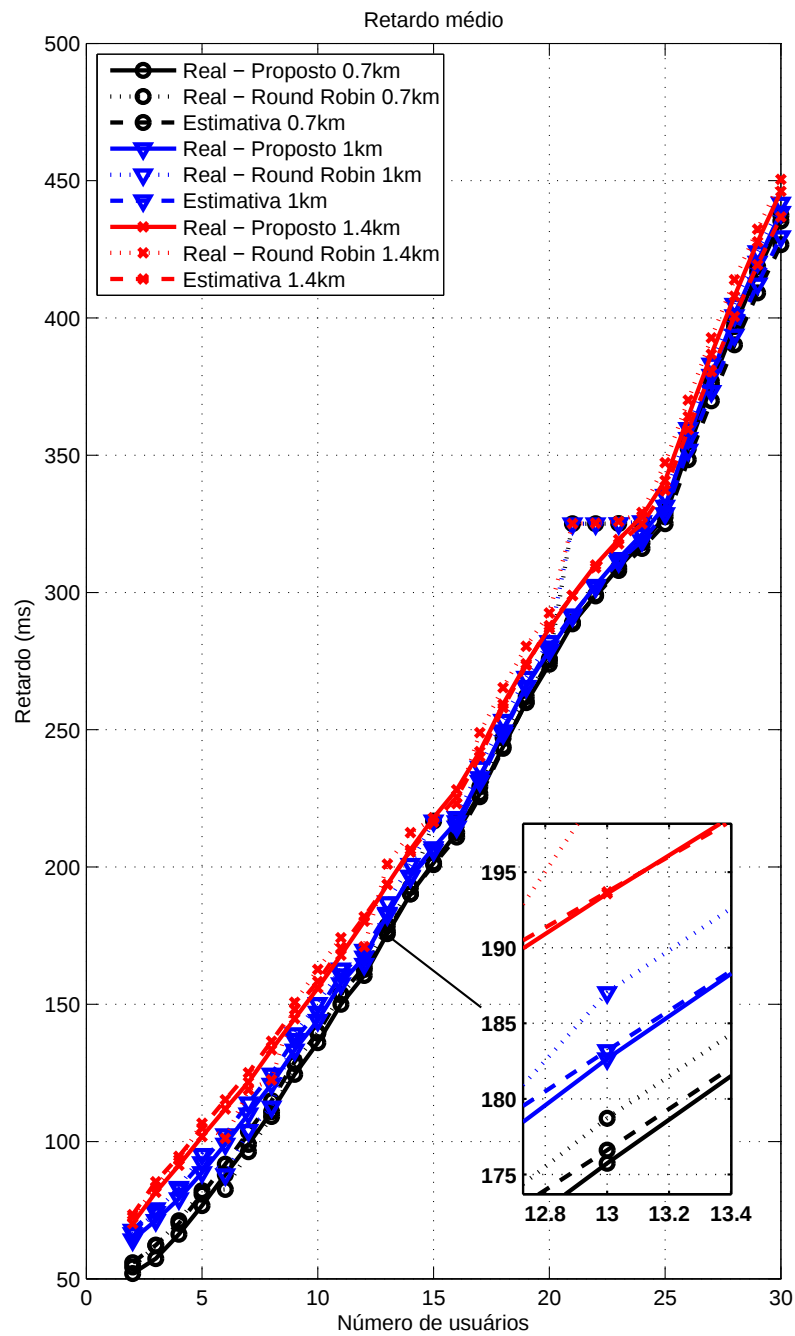


Figura 5.11: Estimativa de limitante de retardo e retardo real médio entre usuários calculado para algoritmo proposto e escalonador Round-Robin, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

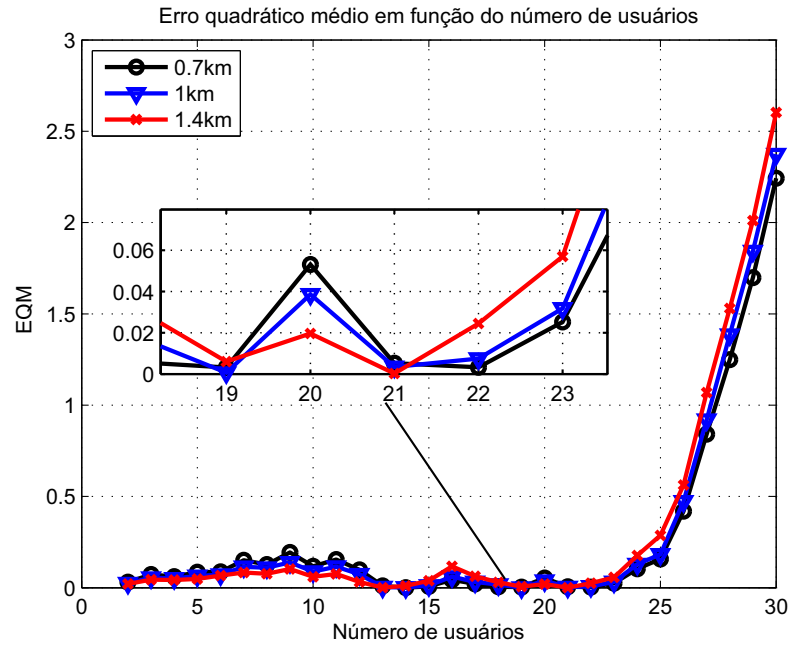


Figura 5.12: Erro Quadrático Médio (EQM) na estimativa de limitante de retardo em função do número de usuários, considerando distâncias de 0.7km, 1km e 1.4km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

5.2 Algoritmo Proposto de Alocação Utilizando Estimação de Retardo

No Capítulo 4, é proposto um algoritmo de escalonamento que procura reduzir o retardo em redes LTE. A principal diferença do algoritmo apresentado no Capítulo 4 e do algoritmo proposto neste capítulo é que neste último o limitante de retardo é estimado a medida que as características dos dados de tráfego no sistema variam, utilizando conceito de Cálculo de Rede, conforme explicado na Seção 5.1. No algoritmo apresentado no Capítulo 4 o retardo foi calculado com base em valores reais passados de retardo da rede. Sendo assim, a vantagem da proposta do presente capítulo é que o algoritmo pode tomar decisões antecipadas utilizando as estimativas e previsões de retardo baseados em uma modelagem do tráfego e do sistema.

O algoritmo proposto, descrito em detalhes no Algoritmo 5.1, pode ser resumido em três fases:

1. Estima-se o número de SBs requeridos para cada usuário com prioridade baseada no ganho médio de canal;

2. Alocam-se os SBs para os usuários de acordo com a prioridade, o limitante de retardo estimado, calculado conforme equação 5-16, e o critério de retardo máximo, definido conforme Tabela 4.2;
3. Alocam-se os SBs remanescentes para os usuários com prioridade definida de acordo com o retardo estimado.

A prioridade na alocação dos recursos é definida em ordem crescente pelo ganho médio do canal por usuário, ou seja, os usuários com piores condições de canal tem maior prioridade. O ganho médio do canal G_k por usuário k é calculado pela seguinte equação:

$$G_k = \frac{1}{N} \sum_{n=1}^N g_{k,n}, \quad (5-17)$$

onde $g_{k,n}$ é o ganho médio do canal para o usuário k no n -ésimo SB e N é o número de SBs disponíveis para *downlink* LTE.

A quantidade N_k de SBs requeridas para cada usuário k é calculada da seguinte forma, com base nas condições do canal:

$$N_k = \text{ceil} \left(\left(\frac{G_k}{G_1 + G_2 + \dots + G_k} \right) * N \right), \quad (5-18)$$

onde G_k é o ganho médio do canal por usuário k , N é o número de SBs disponíveis para *downlink* LTE e $\text{ceil}(x)$ é uma função de arredondamento para o menor inteiro maior que x .

Após calculado a prioridade de alocação, os SBs com maior CQI são alocados de acordo com a quantidade de SBs estimadas para cada usuário. Depois é estimado o limitante de retardo conforme equação 5-16 e verificado se o critério de retardo, definido conforme Tabela 4.2, é satisfeito. Se o critério não é satisfeito o algoritmo continua alocando SBs com maior CQI até satisfazer o critério.

O algoritmo garante a alocação dos SBs de forma justa, uma vez que prioriza os usuários com piores condições do canal com objetivo de satisfazer o critério de retardo, ao mesmo tempo que aloca uma quantidade maior de SBs para os usuários com melhores condições de canal.

Após verificado se o critério de retardo máximo foi satisfeito para todos os usuários, o algoritmo aloca os SBs remanescentes, se houver, priorizando o valor de retardo estimado para cada usuário, ou seja, os usuários com maior valor de retardo tem maior prioridade. O objetivo é reduzir o retardo médio depois de satisfeito o critério.

Algoritmo 5.1: Algoritmo de alocação proposto com estimativa de retardo**Entrada:** K : número de usuários N : número de SBs G : CQI dos usuários (matriz $N \times K$) R_k : taxa mínima requerida dos usuários (vetor $1 \times K$)**Saída:** r_k : vetor de taxas alcançadas por usuário k **1 início****2** $W = \{1, 2, \dots, N\}$ **3** $S_k = \{\}, k \in \{1, 2, \dots, K\}$ **4** $r_k = \{\}, k \in \{1, 2, \dots, K\}$ **5** $\hat{d}_k = \{\}, k \in \{1, 2, \dots, K\}$ **6** Calcula G_k conforme equação 5-17**7** Calcula prioridade 1 em ordem crescente das condições do canal (vetor G_k)**8** Calcula N_k conforme equação 5-18**9** Define d_{max} conforme Tabela 4.2**10 para** $k = 1$ até K de acordo com prioridade 1 **faça****11** Aloca SBs com maior CQI para usuário k no vetor S_k de acordo com a quantidade estimada N_k e remove o SB alocado de W **12** Procura o SB com menor CQI alocado em S_k e determina o maior MCS alcançado para usuário k **13** Calcula a taxa r_k alcançada para usuário k em um *subframe* e estima $\hat{d}_k$ conforme equação 5-16**14 enquanto** $\hat{d}_k \geq d_{max}$ **faça****15** Verifica se há SBs disponíveis ($W \neq \emptyset$)**16** Aloca SBs com maior CQI para usuário k no vetor S_k de acordo com a quantidade estimada N_k e remove o SB alocado de W **17** Procura o SB com menor CQI alocado em S_k e determina o maior MCS alcançado para usuário k **18** Calcula a taxa r_k alcançada para usuário k em um *subframe* e estima $\hat{d}_k$ conforme equação 5-16**19 se** $W \neq \emptyset$ (Verifica se há SBs disponíveis)**20 então****21** Calcula prioridade 2 em ordem decrescente do vetor $\hat{d}_k$ **22 para** $k = 1$ até K de acordo com prioridade 2 **faça****23** Verifica se há SBs disponíveis ($W \neq \emptyset$)**24** Aloca SBs com maior CQI para usuário k no vetor S_k de acordo com a quantidade estimada N_k e remove o SB alocado de W **25** Procura o SB com menor CQI alocado em S_k e determina o maior MCS alcançado para usuário k **26** Calcula a taxa r_k alcançada para usuário k em um *subframe* e estima $\hat{d}_k$ conforme equação 5-16

5.3 Simulações e Resultados

Considerando os parâmetros de transmissão *downlink* no sistema LTE e do modelo de canal descritos na Seção 4.2, avalia-se nesta seção os resultados das simulações do algoritmo de alocação proposto, comparando com o algoritmo PSO (*Particle Swarm Optimization*) [42], o algoritmo com garantia de QoS [24] e o algoritmo de redução de retardo proposto no Capítulo 4. Os valores dos parâmetros de QoS apresentados representam valores médios para os TTIs simulados.

As simulações foram realizadas por meio do software MATLAB versão R2014a, utilizando um microcomputador com a seguinte configuração: Processador Intel Core I5-3570 3.40 GHz, 8GB RAM, HD SATA III 7200 RPM, Windows 8 64bits. Optou-se por implementar as funções e rotinas de simulação ao invés de utilizar ferramentas disponíveis de simulação de rede a fim de simular diversos parâmetros de modelagem de canal e do sistema LTE.

5.3.1 Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 1000 TTIs

Nesta subseção foi considerado uma distância entre usuário e estação base de 1km, largura de banda de 5MHz, modelo multipercorso Rayleigh e taxa mínima fixa. Foram simulados cenários com 2 até 20 usuários em 1000 TTIs.

A Tabela 5.1 e a Figura 5.13 apresentam dados estatísticos referentes à modelagem de canal simulada em termos de SINR.

Tabela 5.1: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Estatística	Valor
Média	24.83
Desvio padrão	11.45
Variância	131.15
Valor mínimo	-35.69
Valor máximo	74.57
Índice de dispersão	5.28
Variância normalizada	0.21
Momento de 3ª ordem (simetria)	-192.05
Momento de 4ª ordem (concentração)	53816.30

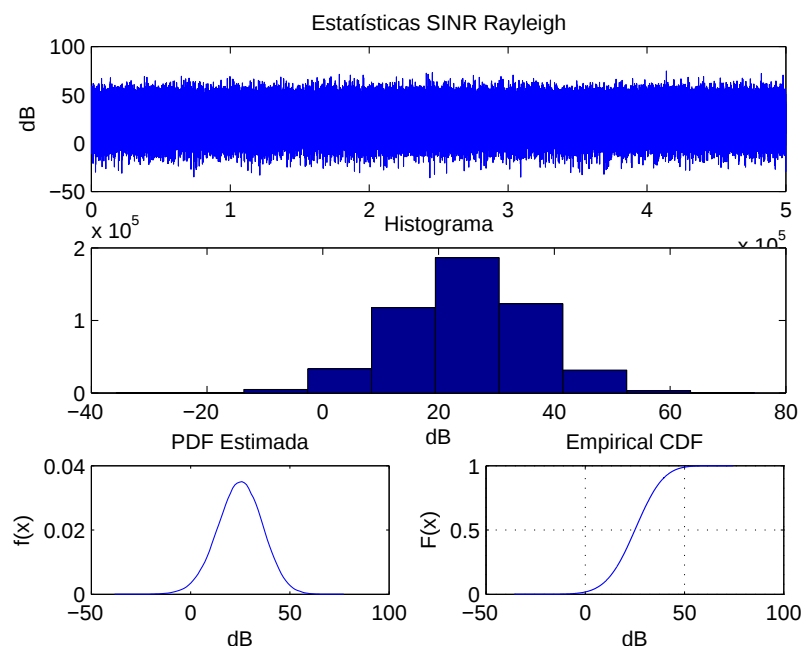


Figura 5.13: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Utilizando a estimativa de limitante de retardo, pode-se verificar na Tabela 5.2 que o algoritmo proposto apresenta uma taxa de atendimento ao critério de retardo máximo superior à taxa apresentada pelo algoritmo de redução de retardo, proposto no

Capítulo 4. Este fato se justifica pois o algoritmo proposto utiliza valores de limitante de retardo estimado ao invés de valores reais, forçando assim o algoritmo a alocar mais recursos para os usuários de forma a atender o requisito de retardo. O algoritmo QoS garantido apresenta a menor taxa de atendimento, com maior número de usuários cujo valor de retardo ultrapassa o critério estabelecido.

Em relação à vazão total, o algoritmo com estimativa de limitante de retardo apresenta melhor desempenho que o algoritmo de redução de retardo, independente do número de usuário considerados em simulação. Nota-se que para um número de usuários menor que 8, a diferença em termos de vazão do algoritmo proposto com estimativa e sem estimativa tende a ser maior, justamente nos pontos onde a estimativa de limitante apresenta valores superiores aos valores reais. Ao mesmo tempo, o desempenho do algoritmo com estimativa tende a melhorar em relação aos demais algoritmos à medida que o número de usuários em simulação aumenta, conforme pode ser verificado na Figura 5.14.

O algoritmo PSO apresenta melhor desempenho em termos de vazão média por usuário, ilustrada na Figura 5.15, se considerado um número menor de usuários. Para mais de 12 usuários, o algoritmo PSO tende a apresentar valores decrescentes devido à sua função de penalidade, enquanto os demais algoritmos estudados tendem a apresentar desempenho similar.

A Figura 5.16 mostra que o algoritmo com estimativa de retardo e o algoritmo de redução de retardo apresentam valores semelhantes em termos de retardo médio entre usuários. Ambos não penalizam os usuários com retardo tendendo ao infinito, ou seja, espera infinita no *buffer*, em virtude de sua característica de balanceamento de usuários. O algoritmo PSO apresenta em geral o pior desempenho em termos de retardo médio.

O algoritmo com estimativa de retardo apresenta valores de índice de justiça similares aos valores apresentados pelo algoritmo de redução de retardo se considerado até 18 usuários em simulação, conforme verificado na Figura 5.17.

A Figura 5.18 mostra que o algoritmo com estimativa de limitante de retardo tende a melhorar o desempenho em relação à perda média do sistema para mais de 18 usuários, em comparação ao algoritmo de redução de retardo. O algoritmo QoS garantido apresenta em geral maiores valores de perda média devido à quantidade de usuários a qual não foram atendidos em transmissão.

A Tabela 5.3 mostra que os algoritmos de redução de retardo e QoS garantido apresentam menores valores de tempo de processamento. O algoritmo PSO apresenta os maiores valores em todos os cenários.

Tabela 5.2: Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	56918	72.76
PSO	14513	93.05
Proposto com estimativa de retardo	2722	98.70
Proposto	3654	98.25

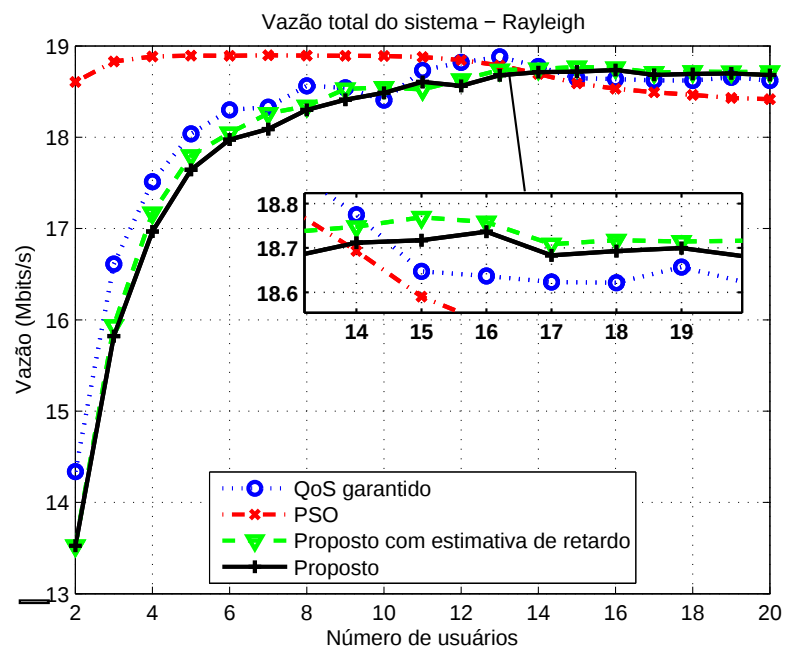


Figura 5.14: Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

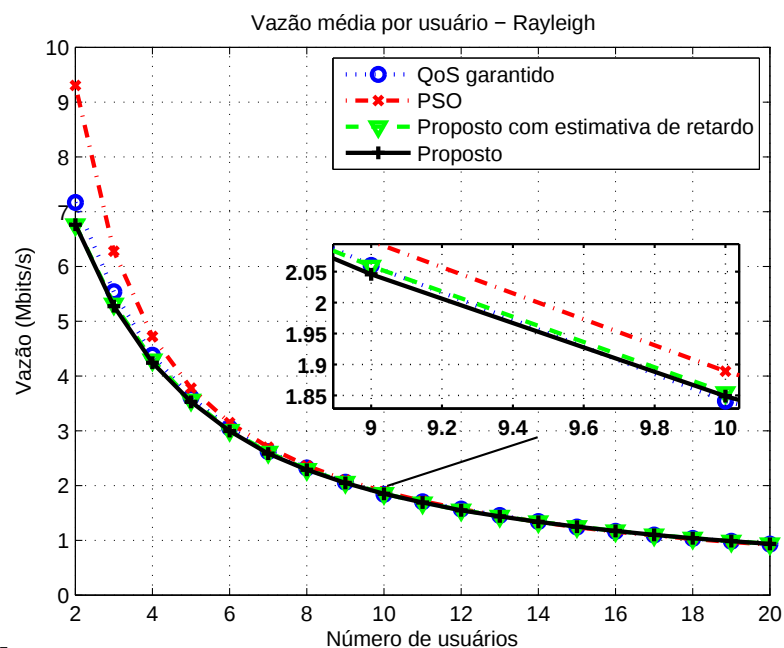


Figura 5.15: Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

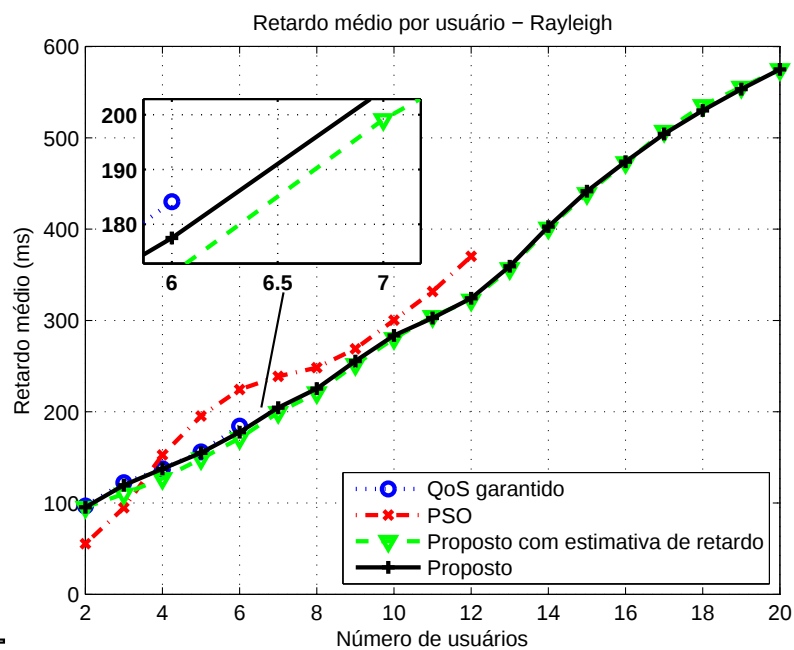


Figura 5.16: Retardo médio entre usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

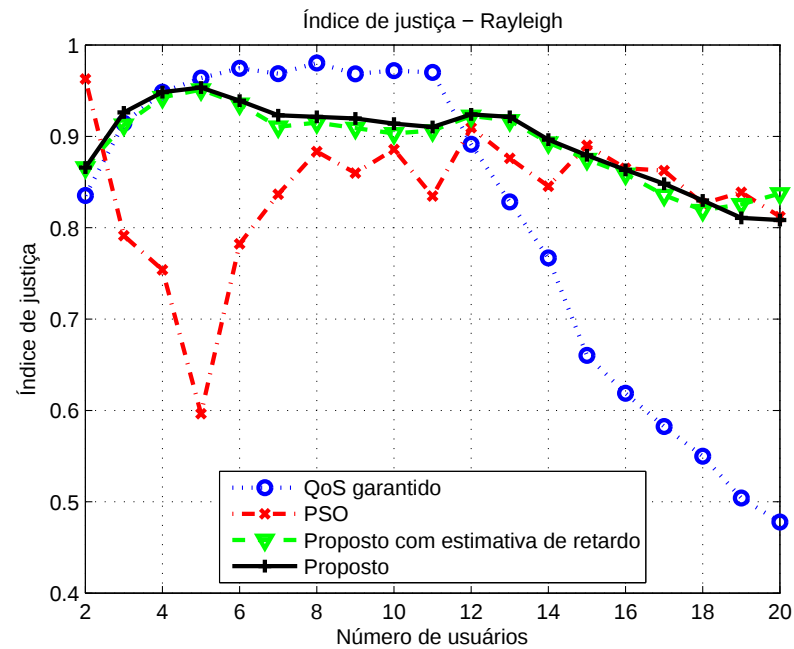


Figura 5.17: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

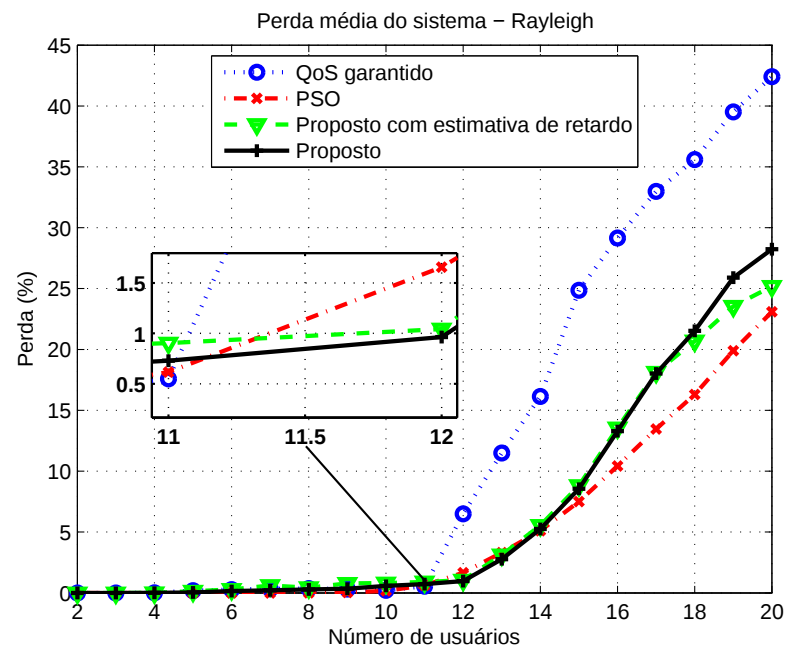


Figura 5.18: Perda média do sistema em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 5.3: *Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs*

Algoritmo	Tempo de processamento (ms)				
	4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS garantido	0.259	0.364	0.467	0.575	0.629
PSO	70.543	71.961	73.535	74.766	75.986
Proposto com estimativa de retardo	0.825	1.523	2.339	2.555	3.986
Proposto	0.300	0.425	0.540	0.565	0.775

5.3.2 Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Fixa e 24 Horas de Simulação

Nesta subseção considerou-se uma distância entre usuário e estação base de 1km, largura de banda de 5MHz, taxa mínima fixa e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 20 usuários em 864×10^5 TTIs, correspondente a um período de 24 horas.

A Tabela 5.4 mostra que o algoritmo proposto com estimativa de retardo apresenta maior taxa de atendimento ao critério, enquanto o algoritmo QoS garantido apresenta a menor taxa de atendimento.

Observa-se nas Figuras 5.19 e 5.20 que o algoritmo proposto com estimativa melhora o desempenho em termos de vazão se comparado ao algoritmo proposto. O algoritmo PSO apresenta em geral maiores valores de vazão.

Em termos de retardo médio entre usuários, o algoritmo proposto com estimativa de retardo apresenta os menores valores em comparação aos demais algoritmos, conforme verificado na Figura 5.21.

Os demais parâmetros de QoS mostrados nas Figuras 5.22 e 5.23 apresentam valores similares aos valores apresentados na subseção anterior, com 1000 TTIs. O fato se repete em relação ao tempo de processamento, mostrado na Tabela 5.5.

Os resultados mostrados nesta subseção demonstram que o cenário com 1000 TTIs, mostrado na subseção anterior, é similar ao cenário com 864×10^5 TTIs.

Tabela 5.4: Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

Algoritmo	Taxa (%)
QoS garantido	73.37
PSO	93.10
Proposto com estimativa de retardo	99.15
Proposto	98.79

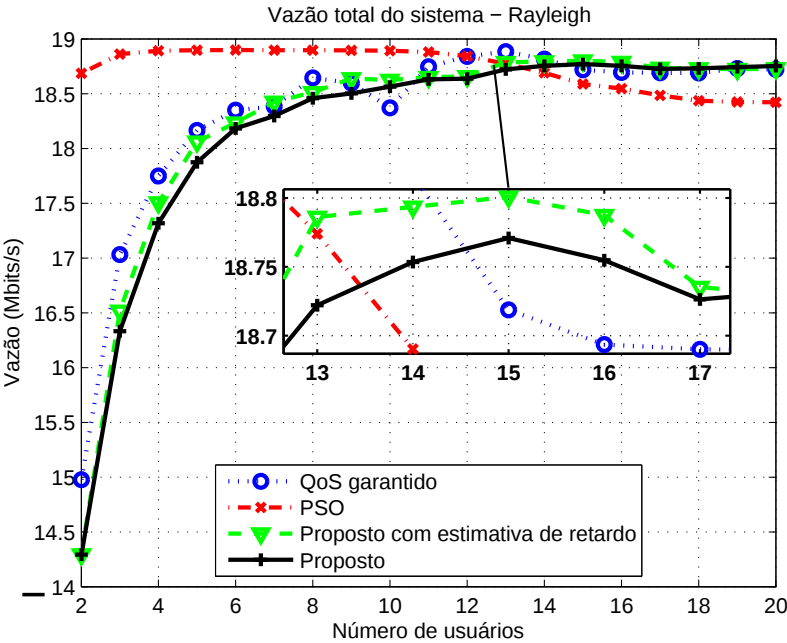


Figura 5.19: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

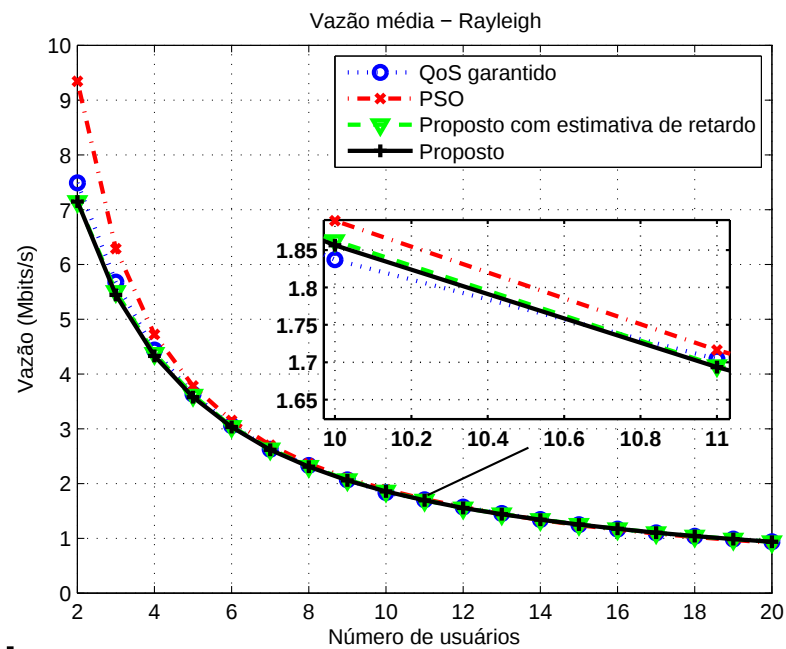


Figura 5.20: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

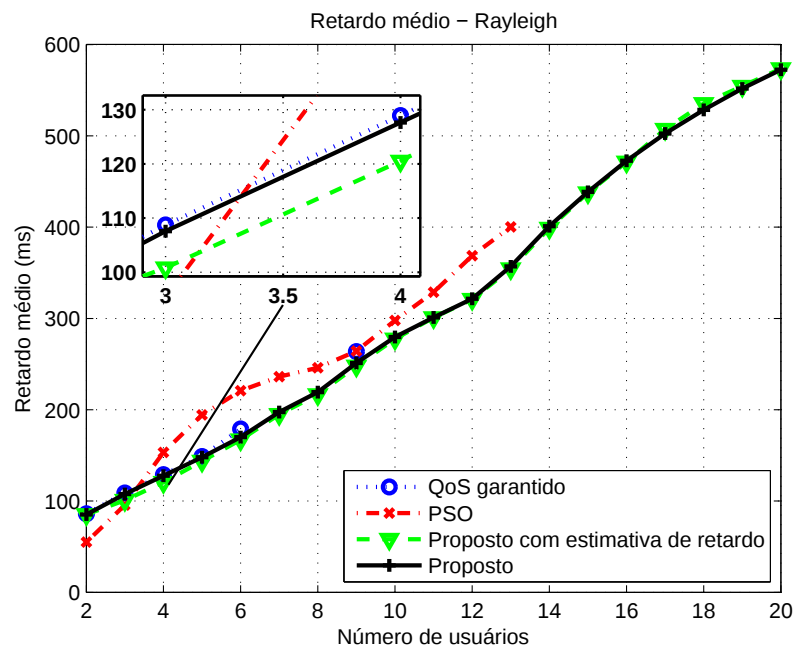


Figura 5.21: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

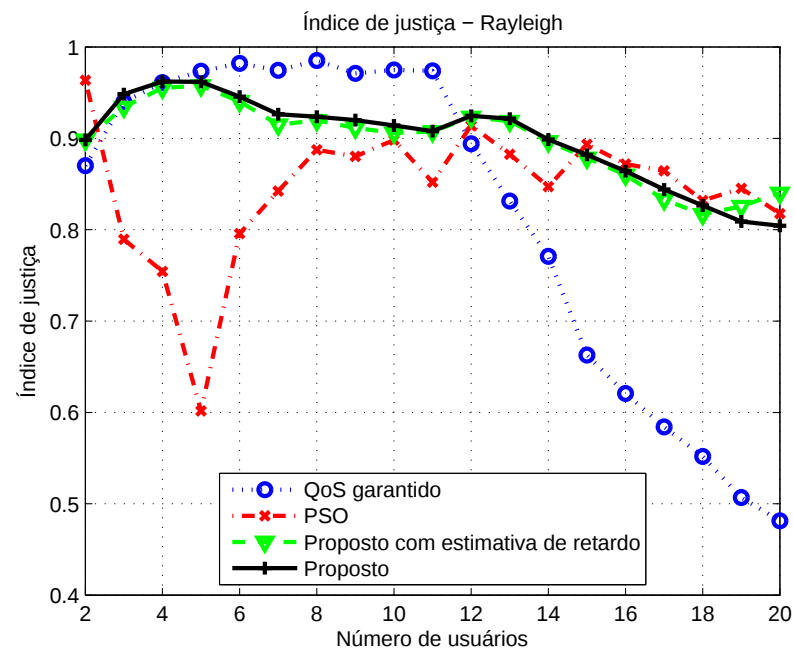


Figura 5.22: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

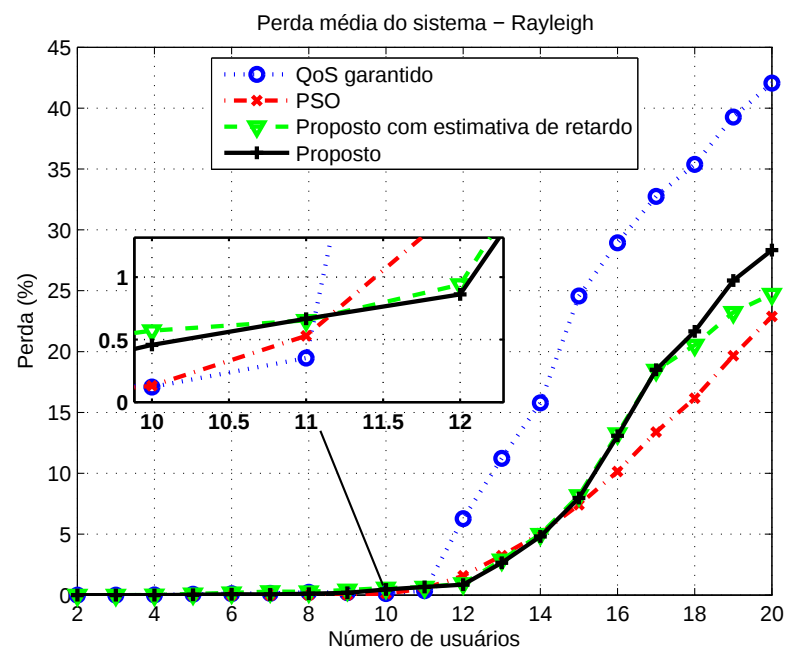


Figura 5.23: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação

Tabela 5.5: *Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 24 horas de simulação*

Algoritmo	Tempo de processamento (ms)				
	4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS garantido	0.258	0.364	0.471	0.575	0.631
PSO	72.338	73.698	75.498	76.638	78.296
Proposto com estimativa de retardo	0.753	1.620	2.508	3.340	4.316
Proposto	0.278	0.402	0.521	0.551	0.764

5.3.3 Distância de 1km, Largura de Banda de 5MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs

Nesta subseção foi considerado uma distância entre usuário e estação base de 1km, largura de banda de 5MHz, modelo multipercurso Rayleigh e taxa mínima calculada adaptativamente conforme explicado na Seção 3.5. Foram simulados cenários com 2 até 20 usuários em 1000 TTIs.

A Tabela 5.6 mostra que o algoritmo proposto com estimativa de retardo possui a maior taxa de atendimento ao critério, com valor próximo ao algoritmo proposto de redução de retardo. Essa característica se justifica pelo fato do algoritmo com estimativa utilizar valores limitantes de maneira a garantir QoS em termos de retardo. O algoritmo PSO apresenta a menor taxa de atendimento.

Em termos de vazão total e vazão média, ilustradas respectivamente nas Figuras 5.24 e 5.25, o algoritmo proposto com estimativa de retardo melhora o desempenho em comparação ao algoritmo de redução de retardo, e apresenta valores similares ao algoritmo QoS garantido para mais de 9 usuários. Em relação à vazão, o algoritmo PSO apresenta melhor desempenho em geral, pois os valores de taxa mínima calculados adaptativamente acarretam em função de penalidade mínima.

O algoritmo proposto com estimativa de retardo apresenta pequena melhora no desempenho em termos de retardo médio por usuário quando comparado ao algoritmo de redução de retardo, conforme pode ser verificado na Figura 5.26, pois os valores de limitante estimados são em média superiores aos valores reais e possibilitam ao algoritmo fornecer blocos de forma a garantir o atendimento a este quesito.

O algoritmo QoS garantido apresenta os maiores valores de índice de justiça, conforme pode ser visto na Figura 5.27. O algoritmo proposto com estimativa de retardo tem desempenho inferior ao algoritmo proposto, e o algoritmo PSO apresenta o pior desempenho dentre os algoritmos em termos do critério de justiça.

Em termos de perda média do sistema, ilustrada na Figura 5.28, os algoritmos propostos apresentam o melhor desempenho no geral, principalmente se considerado mais de 16 usuários em simulação. A característica de balanceamento dos algoritmos propostos garante que todos os usuários sejam atendidos e consequentemente ocorra menos perda.

Assim como ocorre na subseção anterior, o algoritmo proposto com estimativa de retardo apresenta valores de tempo de processamento ligeiramente superiores aos valores apresentados pelos algoritmo proposto e pelo algoritmo QoS garantido. O algoritmo PSO apresenta o pior desempenho, conforme verificado na Tabela 5.7.

Tabela 5.6: *Taxa de atendimento ao critério de retardo máximo para total de 209000 alocações, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs*

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	18888	90.96
PSO	28905	86.17
Proposto com estimativa de retardo	2722	98.70
Proposto	4287	97.95

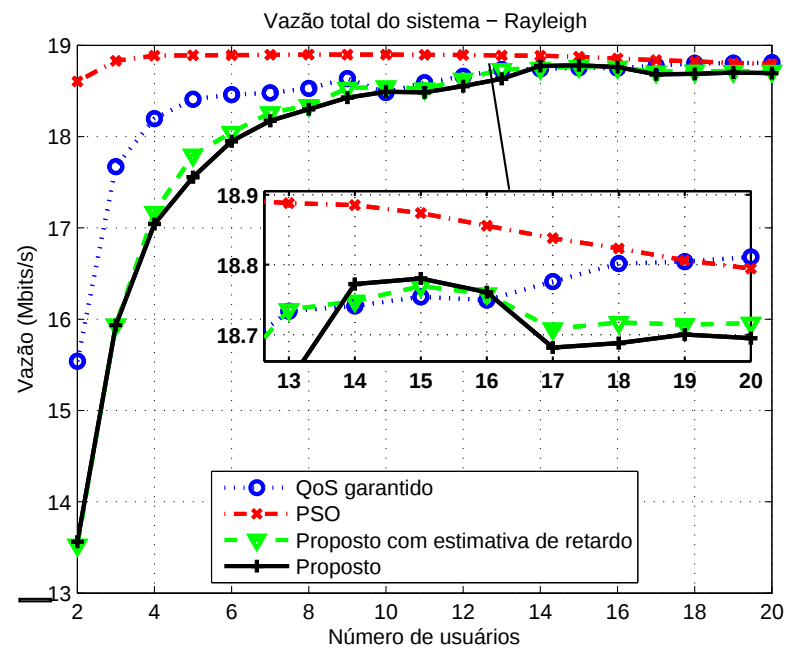


Figura 5.24: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

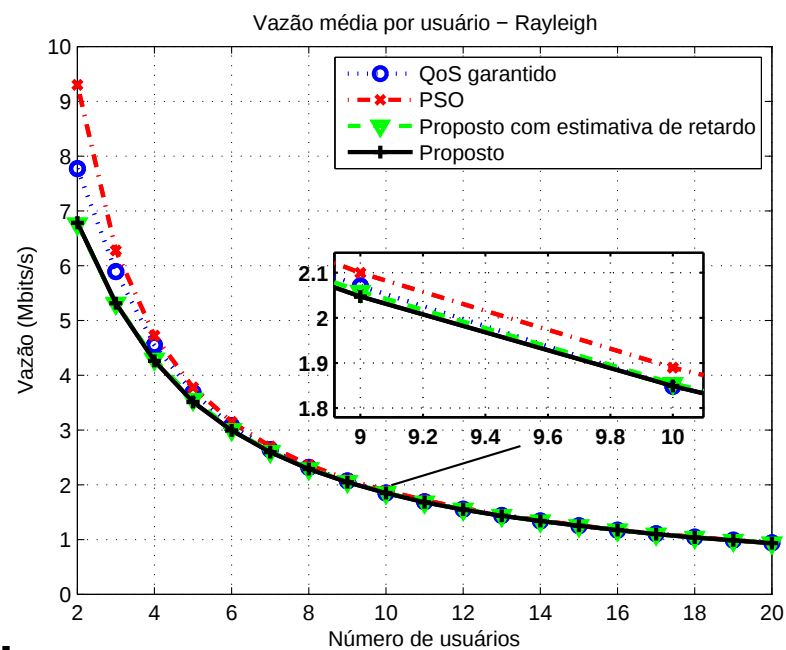


Figura 5.25: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

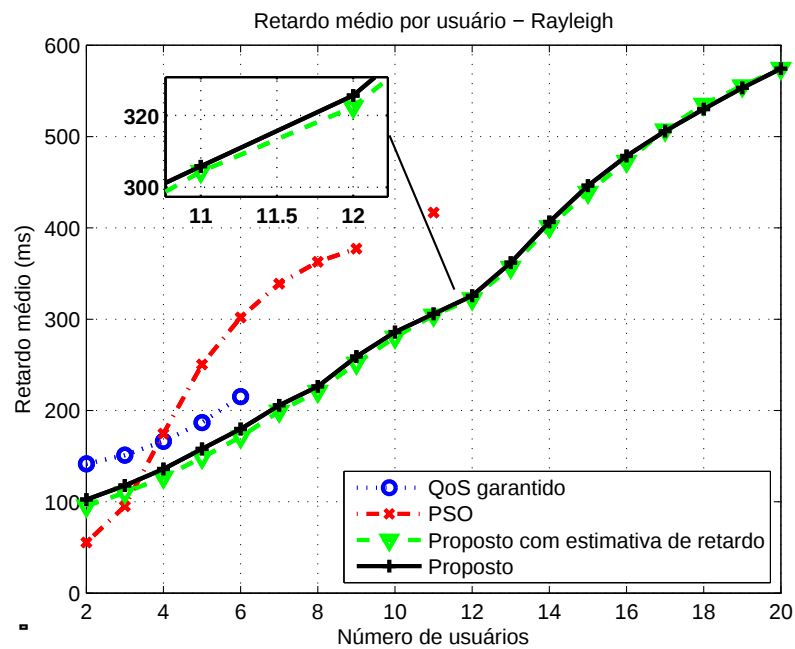


Figura 5.26: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

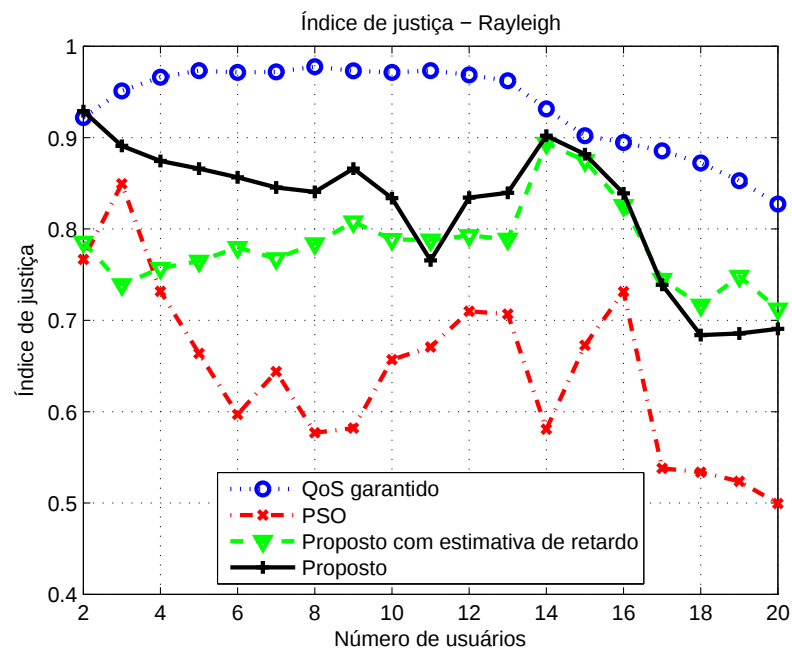


Figura 5.27: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

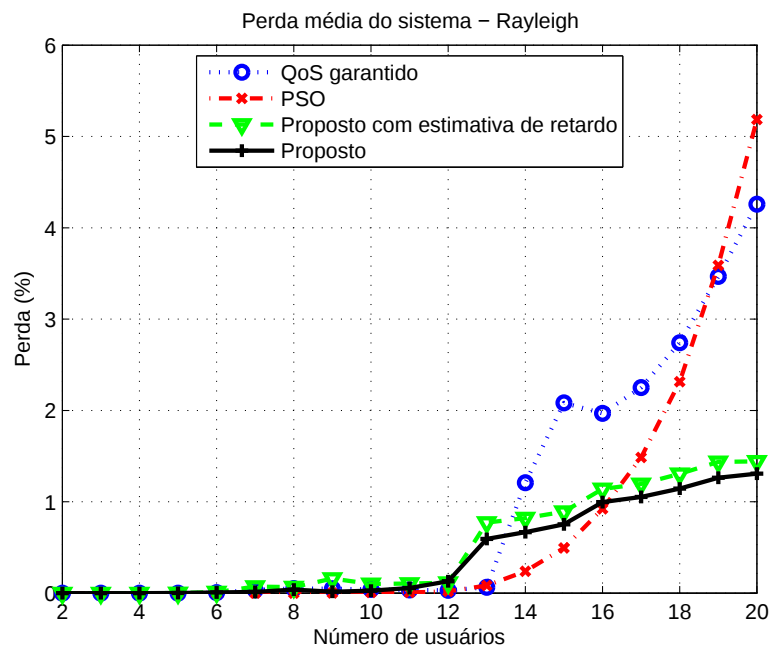


Figura 5.28: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 5.7: Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Tempo de processamento (ms)				
	4 usuários	8 usuários	12 usuários	16 usuários	20 usuários
QoS garantido	0.250	0.350	0.442	0.536	0.621
PSO	70.469	72.087	73.495	74.736	76.037
Proposto com estimativa de retardo	0.771	1.457	2.276	2.473	3.898
Proposto	0.268	0.392	0.500	0.529	0.733

5.3.4 Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 1000 TTIs

Nesta subseção foi considerado uma distância entre usuário e estação base de 1km, largura de banda de 10MHz, modelo multipercurso Rayleigh e taxa mínima fixa. Foram simulados cenários com 2 até 30 usuários em 1000 TTIs.

A Tabela 5.8 e a Figura 5.29 apresentam dados estatísticos referentes à modelagem de canal simulada em termos de SINR.

Tabela 5.8: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Estatística	Valor
Média	21.85
Desvio padrão	11.45
Variância	131.12
Valor mínimo	-46.48
Valor máximo	79.92
Índice de dispersão	5.99
Variância normalizada	0.27
Momento de 3ª ordem (simetria)	-202.84
Momento de 4ª ordem (concentração)	54068.9

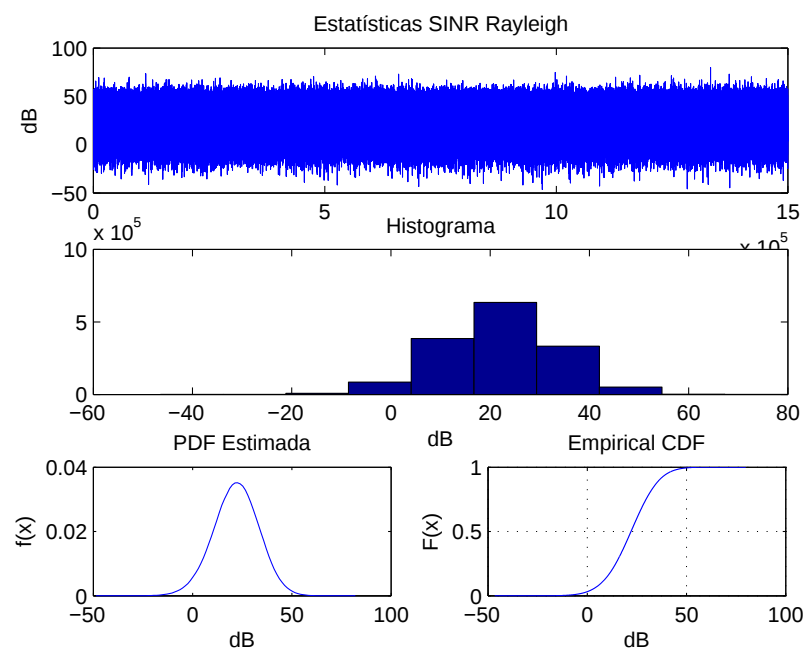


Figura 5.29: Estatísticas da modelagem de canal em termos de SINR medido em dB, considerando distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

A Tabela 5.9 mostra que o algoritmo proposto com estimativa de retardo aumenta a taxa de atendimento ao critério de retardo máximo, comparando com o algoritmo proposto, em virtude de utilizar valores em média superiores aos valores reais se considerado até 13 usuários, o que possibilita ao algoritmo atender aos requisitos de retardo com maior facilidade. O algoritmo QoS garantido apresenta a menor taxa.

Em termos de vazão total e média, o algoritmo proposto com estimativa de retardo mantém desempenho similar ao algoritmo QoS garantido para mais de 13 usuários, e superior ao algoritmo proposto, conforme visto nas Figuras 5.30 e 5.31. O desempenho do algoritmo proposto com estimativa de retardo é inclusive superior ao desempenho do algoritmo PSO para mais de 18 usuários, pois o algoritmo PSO mantém tendência de decrescimento de seus valores devido a sua função de penalidade.

O algoritmo proposto com estimativa de retardo mantém desempenho similar ao algoritmo proposto em relação ao retardo médio por usuário, com pouca variação, e desempenho superior ao algoritmo QoS. A Figura 5.32 mostra que o algoritmo PSO apresenta o pior desempenho.

Em relação ao índice de justiça, o algoritmo proposto com estimativa de retardo e o algoritmo proposto mantém valores similares. Os valores tendem a ser maiores que os valores apresentados pelo algoritmo QoS garantido para mais de 24 usuários, conforme verificado na Figura 5.33. O fato é que a característica de balanceamento dos dois algoritmos propostos promove desempenho satisfatório em termos de critério de justiça independente do número de usuários considerados no cenário simulado, ao contrário do algoritmo QoS garantido.

A Figura 5.34 mostra que os algoritmos propostos apresentam os menores valores para perda média do sistema em comparação aos algoritmos PSO e algoritmo QoS garantido, novamente em virtude de suas características de balanceamento.

A Tabela 5.10 apresenta os valores calculados de tempo de processamento dos algoritmos estudados. Nota-se que o algoritmo proposto e o algoritmo QoS garantido apresentam os menores valores, enquanto o algoritmo PSO apresenta os maiores valores.

Tabela 5.9: Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	46927	89.88
PSO	15943	96.56
Proposto com estimativa de retardo	3445	99.26
Proposto	4761	98.97

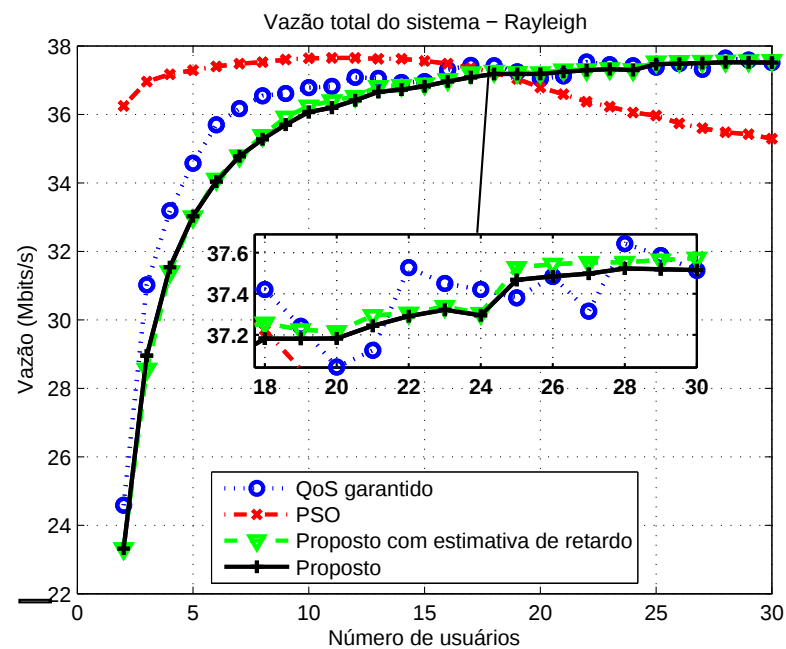


Figura 5.30: Vazão total em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

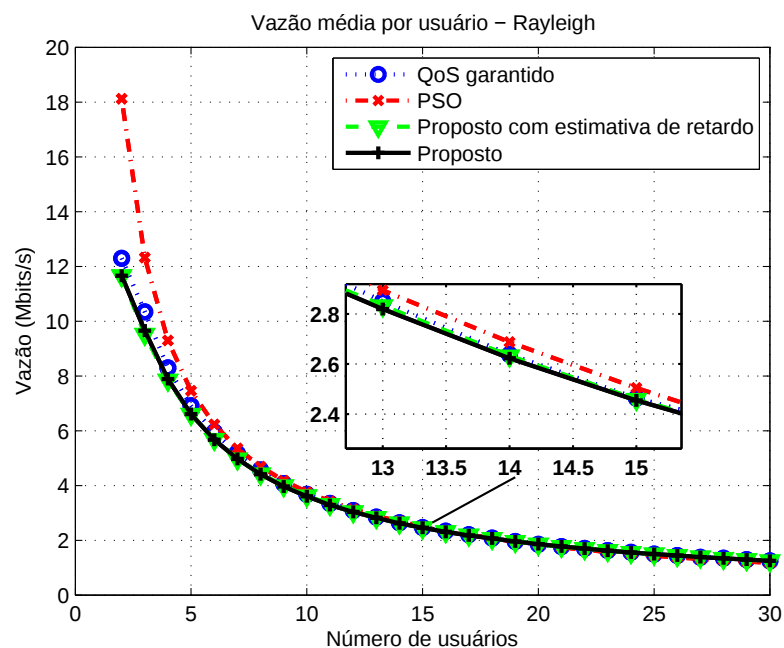


Figura 5.31: Vazão média em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

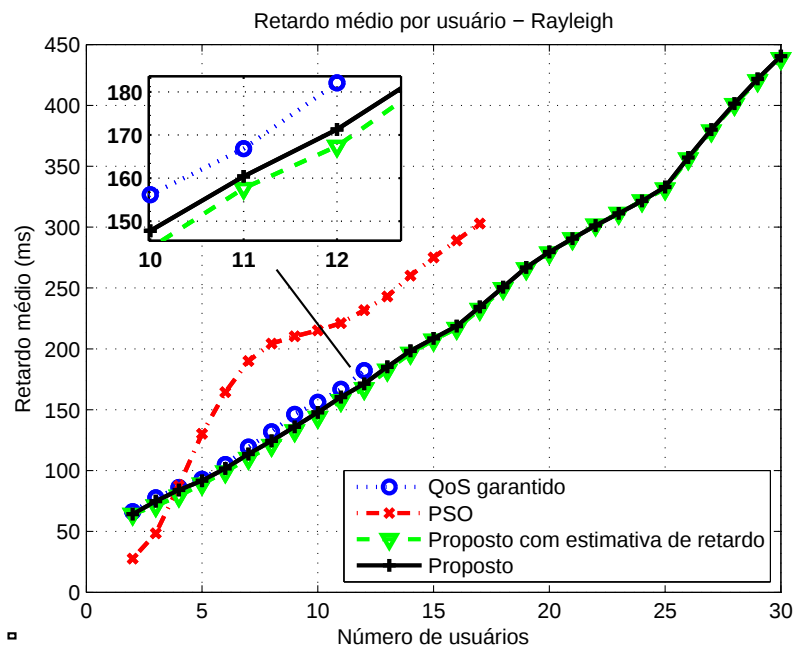


Figura 5.32: Retardo médio entre usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

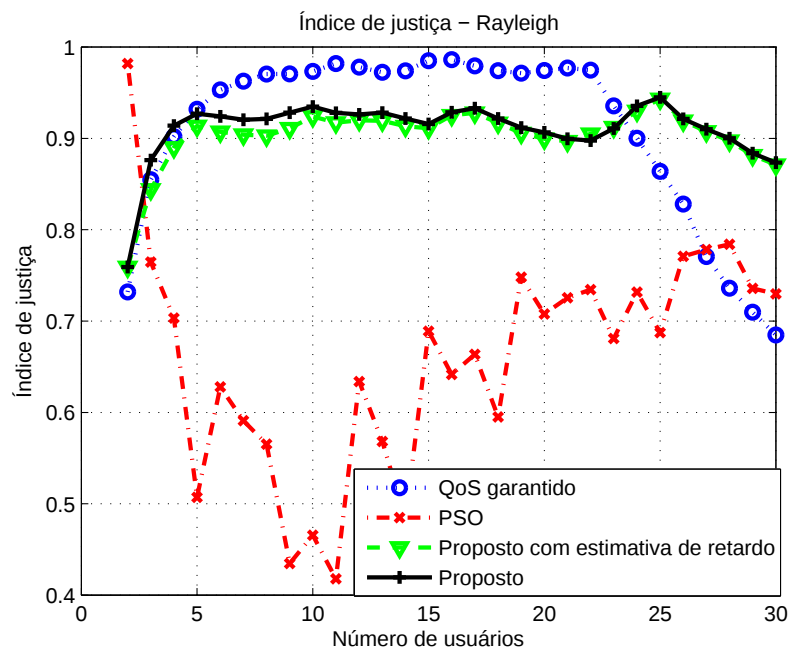


Figura 5.33: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

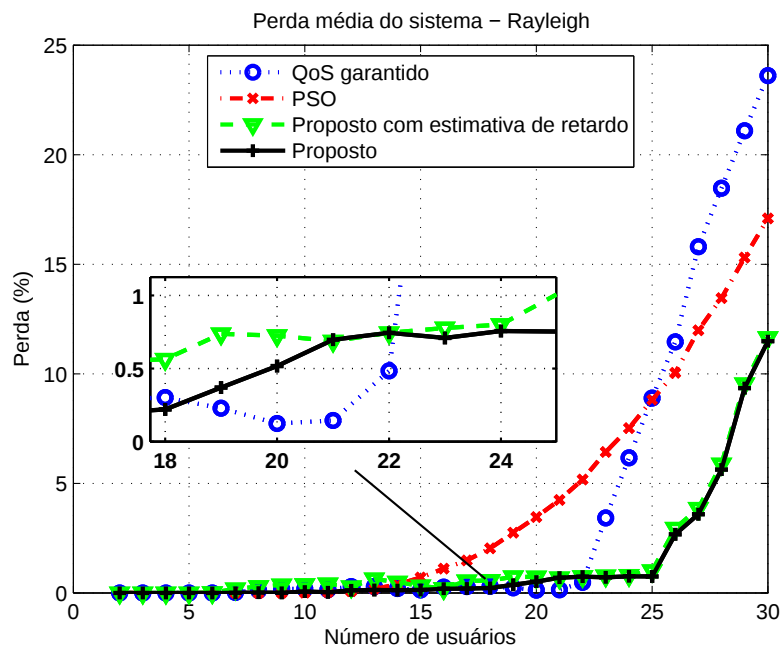


Figura 5.34: Perda média do sistema em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 5.10: *Tempo de processamento em função do número de usuários, considerando taxa mínima fixa, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs*

Algoritmo	Tempo de processamento (ms)				
	6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS garantido	0.423	0.603	0.756	0.878	1.057
PSO	80.801	84.931	88.808	92.238	96.674
Proposto com estimativa de retardo	1.242	2.539	3.029	4.979	4.803
Proposto	0.461	0.665	0.761	0.984	1.050

5.3.5 Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Fixa e 24 Horas de Simulação

Nesta subseção considerou-se uma distância entre usuário e estação base de 1km, largura de banda de 10MHz, taxa mínima fixa e modelo multipercurso Rayleigh. Foram simulados cenários com 2 até 30 usuários em 864×10^5 TTIs, correspondente a um período de 24 horas.

Em relação à taxa de atendimento ao critério de retardo máximo, mostrada na Tabela 5.11, o cenário se mantém o mesmo. O algoritmo proposto com estimativa apresenta a maior taxa de atendimento, enquanto o algoritmo QoS garantido apresenta a menor.

Em termos de vazão total e vazão média, mostradas nas Figuras 5.35 e 5.36, o algoritmo PSO apresenta em geral os maiores valores. O algoritmo proposto com estimativa de retardo apresenta melhora no desempenho da vazão se comparado ao algoritmo proposto, e se equipara ao algoritmo QoS garantido para cenário com mais de 11 usuários.

A Figura 5.37 mostra que o algoritmo proposto e o algoritmo proposto com estimativa apresentam menores valores de retardo médio entre usuários, assim como no cenário com 1000 TTIs.

Em relação ao índice de justiça, mostrado na Figura 5.38, e à taxa de perda média do sistema, mostrada na Figura 5.39, o desempenho é o mesmo se comparado ao cenário com 1000 TTIs. O algoritmo proposto com estimativa e o algoritmo proposto apresentam em geral menores taxas de perda.

Os algoritmos em estudo apresentam valores de tempo de processamento similares aos valores apresentados no cenário com 1000 TTIs, conforme pode ser verificado nas Tabelas 5.12 e 5.10.

Os resultados demonstram que os cenários com 1000 e com 864×10^5 TTIs apresentam desempenho similar neste estudo de caso.

Tabela 5.11: Taxa de atendimento ao critério de retardo máximo, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

Algoritmo	Taxa (%)
QoS garantido	89.51
PSO	85.88
Proposto com estimativa de retardo	98.97
Proposto	98.62

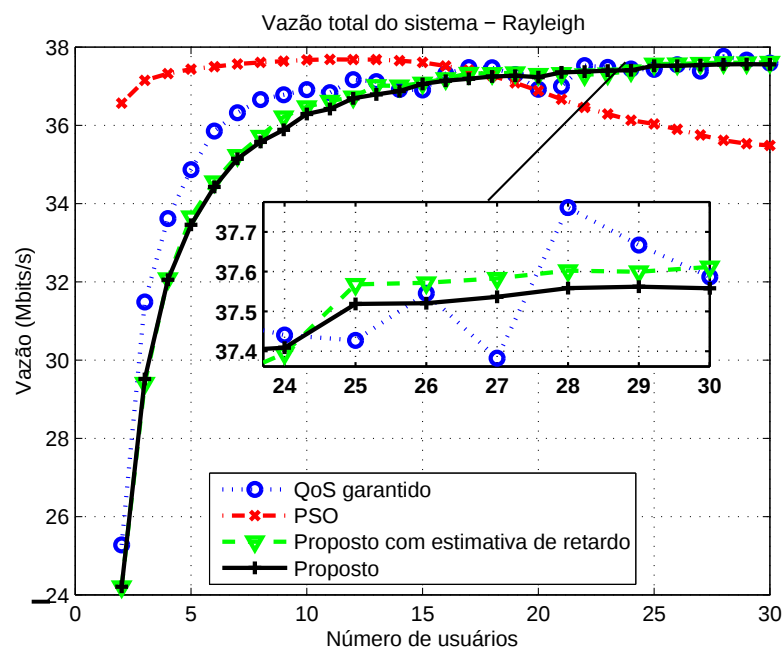


Figura 5.35: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

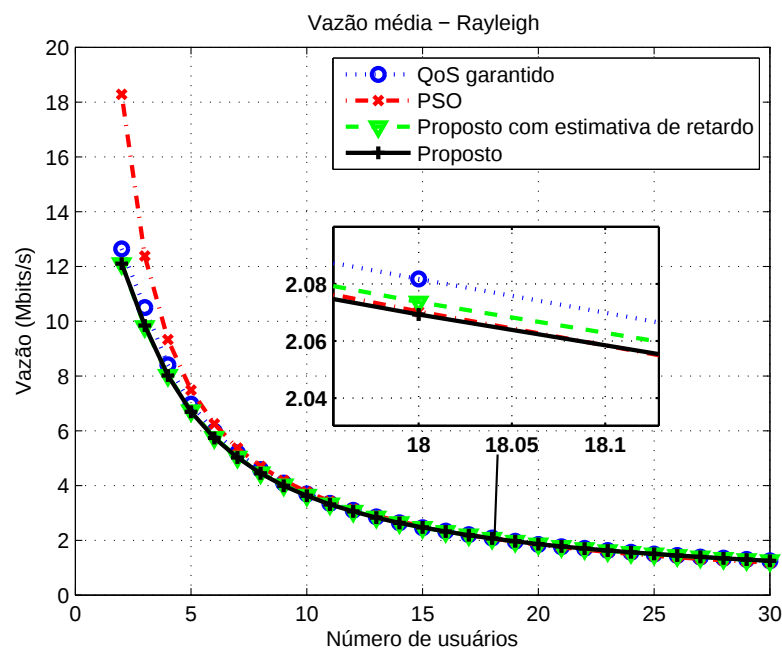


Figura 5.36: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

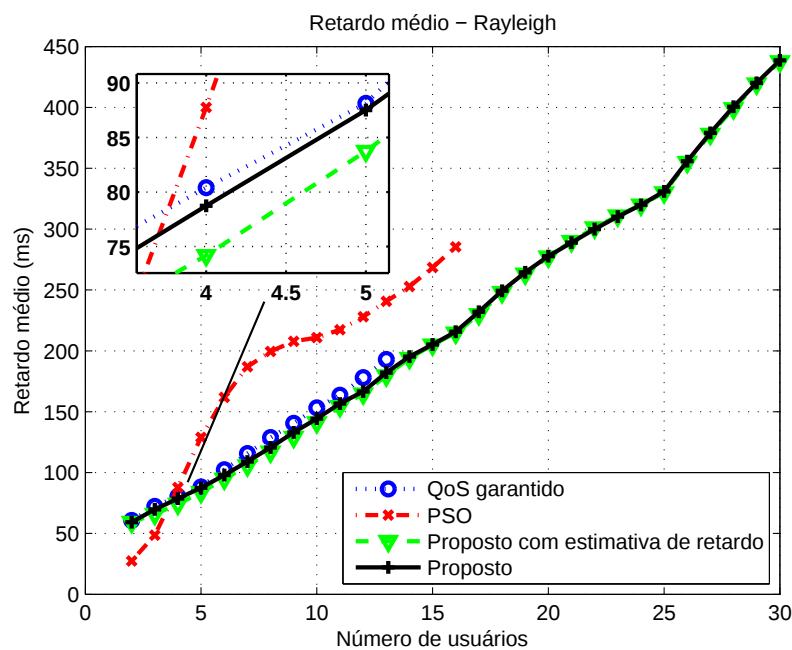


Figura 5.37: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

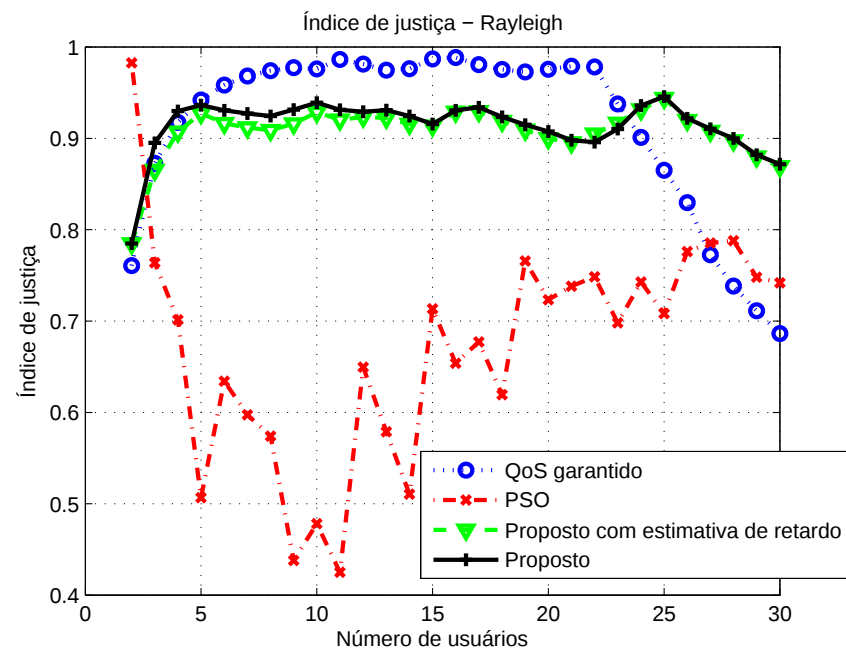


Figura 5.38: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

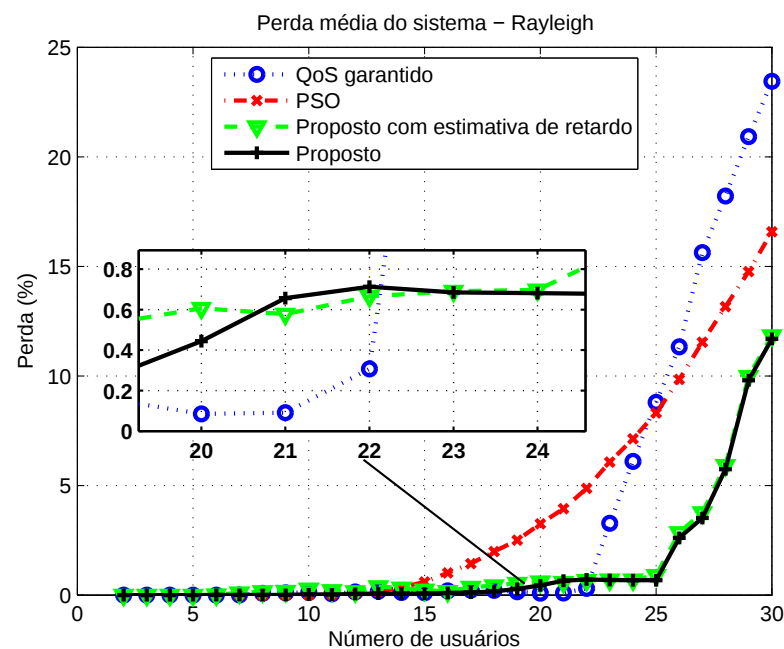


Figura 5.39: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação

Tabela 5.12: *Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 24 horas de simulação*

Algoritmo	Tempo de processamento (ms)				
	6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS garantido	0.430	0.611	0.768	0.894	1.073
PSO	81.020	85.217	88.822	93.149	96.345
Proposto com estimativa de retardo	1.337	2.789	3.194	4.892	5.011
Proposto	0.446	0.645	0.754	0.985	1.053

5.3.6 Distância de 1km, Largura de Banda de 10MHz, Taxa Mínima Calculada Adaptativamente e 1000 TTIs

Nesta subseção foi considerado uma distância entre usuário e estação base de 1km, largura de banda de 10MHz, modelo multipercurso Rayleigh e taxa mínima calculada adaptativamente conforme explicado na Seção 3.5. Foram simulados cenários com 2 até 30 usuários em 1000 TTIs.

Da mesma forma que nos cenários simulados anteriormente, o algoritmo proposto com estimativa de retardo apresenta maior taxa de atendimento ao critério de retardo definido em comparação ao algoritmo proposto, conforme verificado na Tabela 5.13. O algoritmo PSO apresenta a menor taxa.

Em termos de vazão total e vazão média, ilustradas nas Figuras 5.40 e 5.41, o algoritmo PSO apresenta o melhor desempenho em geral. O desempenho do algoritmo proposto com estimativa de retardo é superior ao desempenho do algoritmo proposto, e superior ao QoS garantido se considerado mais de 18 usuários em simulação, devido ao algoritmo proposto estimar número maior de blocos para usuários com melhores condições de canal, garantindo assim melhor desempenho no quesito vazão.

Nota-se na Figura 5.42 que o algoritmo proposto com estimativa de retardo apresenta os menores valores de retardo médio por usuário, similares aos valores apresentados pelo algoritmo proposto. O algoritmo PSO apresenta os maiores valores.

Em relação ao índice de justiça, ilustrado na Figura 5.43, o algoritmo QoS garantido apresenta o melhor desempenho, enquanto o algoritmo PSO apresenta os menores valores. O algoritmo proposto apresenta em geral valores superiores aos valores apresentados pelo algoritmo proposto com estimativa de retardo.

Os algoritmos propostos apresentam valores de taxa de perda média do sistema em geral inferiores aos valores apresentados pelo algoritmo PSO em virtude da características de balanceamento presente nestes, conforme verificado na Figura 5.44.

Quanto ao tempo de processamento, assim como nos cenários simulados anteriormente, o algoritmo QoS garantido e o algoritmo proposto apresentam os menores valores, conforme pode ser verificado na Tabela 5.14.

Tabela 5.13: Taxa de atendimento ao critério de retardo máximo para total de 464000 alocações, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Algoritmo	Quantidade de vezes que usuário não atingiu critério	Taxa (%)
QoS garantido	9667	97.91
PSO	29829	93.57
Proposto com estimativa de retardo	3445	99.26
Proposto	5658	98.78

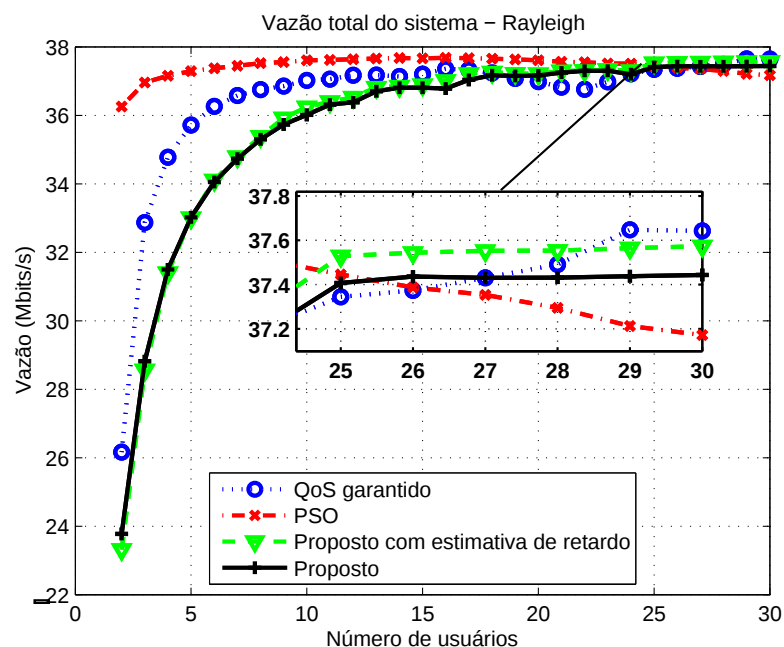


Figura 5.40: Vazão total em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

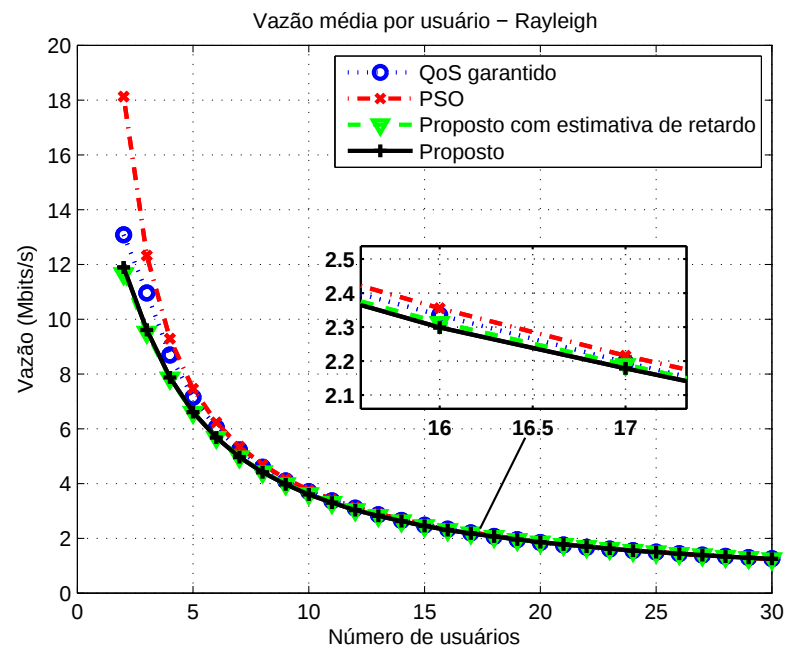


Figura 5.41: Vazão média em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

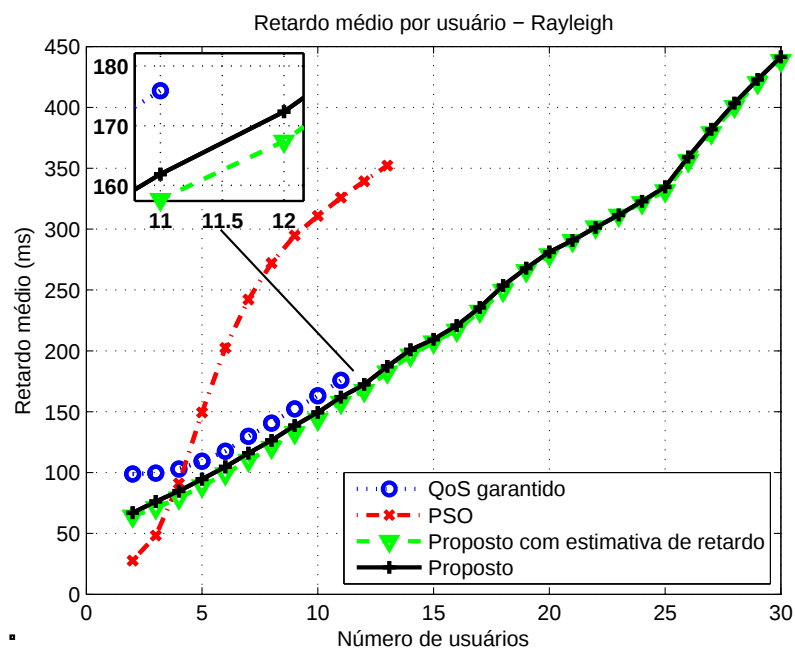


Figura 5.42: Retardo médio entre usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

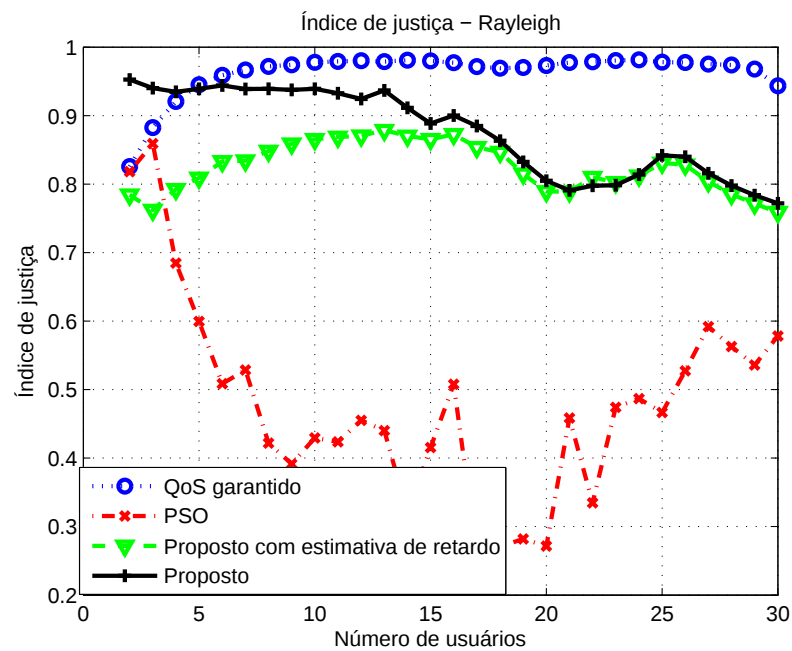


Figura 5.43: Índice de justiça (fairness) em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

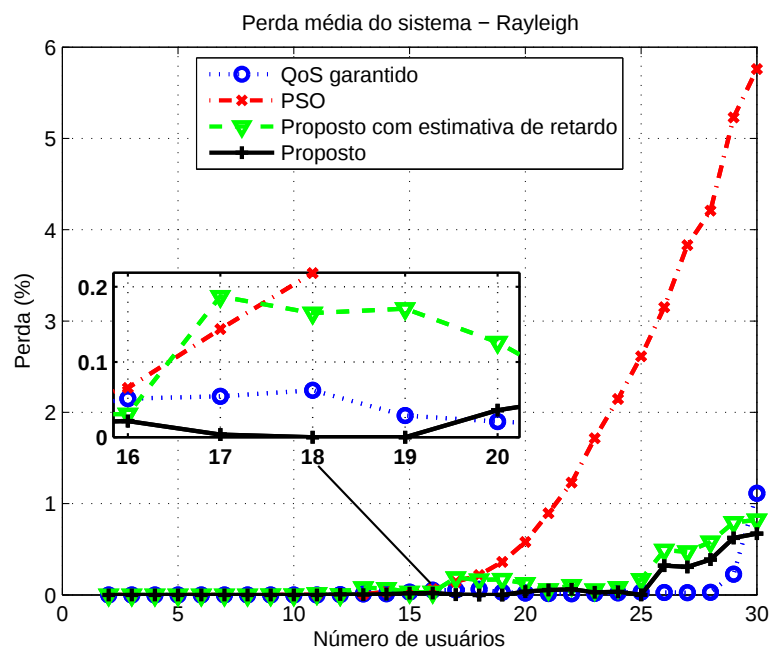


Figura 5.44: Perda média do sistema em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

Tabela 5.14: *Tempo de processamento em função do número de usuários, considerando taxa mínima calculada adaptativamente, distância de 1km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs*

Algoritmo	Tempo de processamento (ms)				
	6 usuários	12 usuários	18 usuários	24 usuários	30 usuários
QoS garantido	0.436	0.605	0.757	0.894	1.064
PSO	78.632	82.739	86.816	90.510	94.081
Proposto com estimativa de retardo	1.266	2.572	3.076	5.058	4.868
Proposto	0.459	0.667	0.767	0.989	1.048

5.3.7 Controle de Admissão Aplicado em Transmissões em Tempo Real

A União Internacional das Telecomunicações (ITU, *International Telecommunication Union*) orienta a respeito do retardo na rede para aplicações de voz em tempo real no documento G.114 [27]. Este documento define três faixas de retardo, como mostra a Tabela 5.15.

Tabela 5.15: *Especificações de retardo recomendadas [10]*

Intervalo (ms)	Descrição
0-150	Aceitável para a maioria das aplicações
150-400	Aceitável, desde que os administradores da rede estejam cientes do tempo de transmissão e do impacto que este tem sobre a qualidade nas aplicações
Acima de 400	Inaceitável para planejamento de redes em geral. No entanto, reconhece-se que em alguns casos excepcionais este limite é excedido

Em geral, o valor de retardo de 200ms é considerado razoável para transmissões de voz em tempo real, e o valor de 250ms pode ser considerado como valor limite. As redes devem ser projetadas de forma que o retardo máximo esperado para aplicações de voz seja conhecido e minimizado [10].

As Figuras 5.45, 5.46 e 5.47 apresentam o retardo médio por usuário e a estimativa do limitante de retardo calculado nas simulações realizadas com o algoritmo proposto neste capítulo, além de definir o retardo máximo tolerável em 250ms, conforme recomendações da ITU [27].

Considerando largura de banda de 5MHz, verifica-se nos cenários simulados que o algoritmo proposto neste capítulo garante o retardo máximo tolerável de 250ms para 9 usuários na alocação à uma distância entre usuário e estação base de 0.7 e 1km, conforme pode ser visto nas Figuras 5.45 e 5.46. Se considerado uma distância de 1.4km, o algoritmo proposto garante o retardo máximo para até 8 usuários, conforme pode ser visto na Figura 5.47.

As Figuras 5.48, 5.49 e 5.50 apresentam o cenário com largura de banda de 10MHz. O algoritmo proposto garante em média o retardo máximo tolerável de 250ms para alocação com 18 usuários nos cenários com distância de 0.7 e 1km. Para uma distância de 1.4km o algoritmo garante o retardo máximo para alocação com 17 usuários.

Verifica-se que a estimativa de limitante de retardo apresenta valores superiores aos valores reais calculados em rede se considerado até 6 usuários para largura de banda de 5MHz e até 13 usuários para largura de banda de 10MHz. Nestes casos, a estimativa de limitante de retardo fornece com precisão valores adequados de maneira a garantir o atendimento ao requisito de retardo do fluxo de tráfego.

Em geral, os valores de retardo estimados se mantêm bastante próximos aos valores calculados. Dessa forma, antes mesmo de se iniciar o escalonamento dos recursos disponíveis, a estimativa pode fornecer o número de usuários que podem ser admitidos para atender um certo valor máximo de retardo. Ou seja, pode-se utilizar o algoritmo proposto de estimação de retardo para se estimar quantos usuários podem ser aceitos de modo que o retardo médio fique abaixo de um valor tolerável. Pode-se utilizar a estimativa de limitante de retardo a fim de fornecer subsídios para que o algoritmo de alocação possa admitir 9 usuários, por exemplo, de maneira a garantir em média o retardo máximo de 250ms considerando largura de banda de 5MHz, ou por exemplo, 18 usuários, no caso da largura de banda de 10MHz.

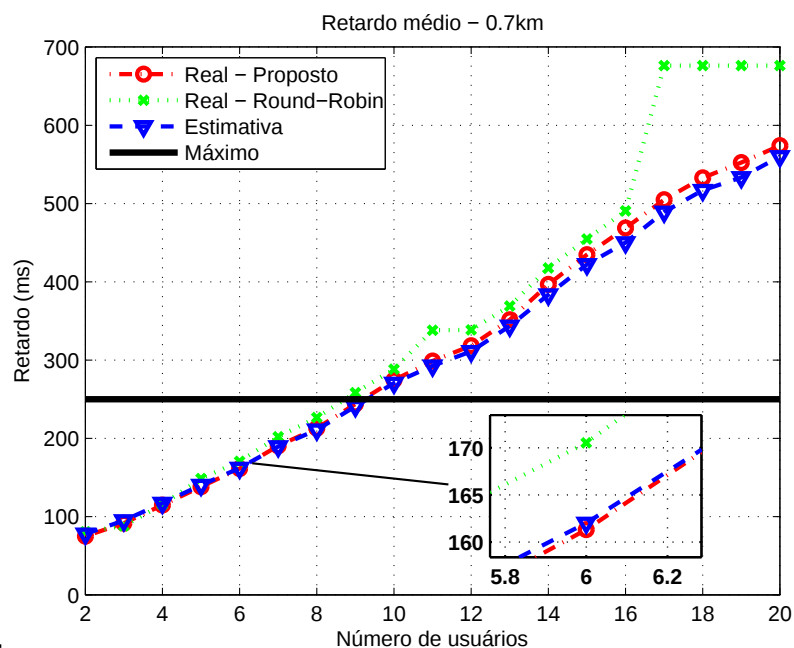


Figura 5.45: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 0.7km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

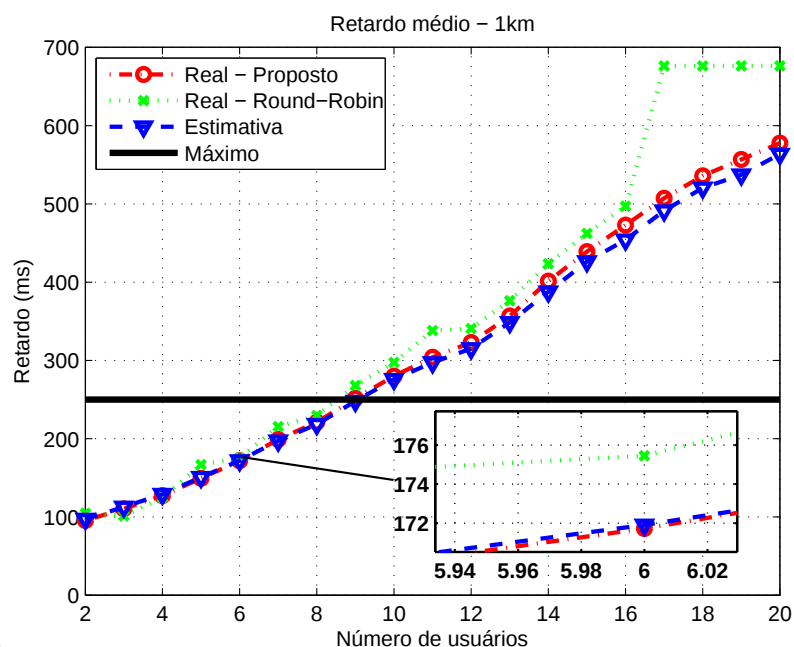


Figura 5.46: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

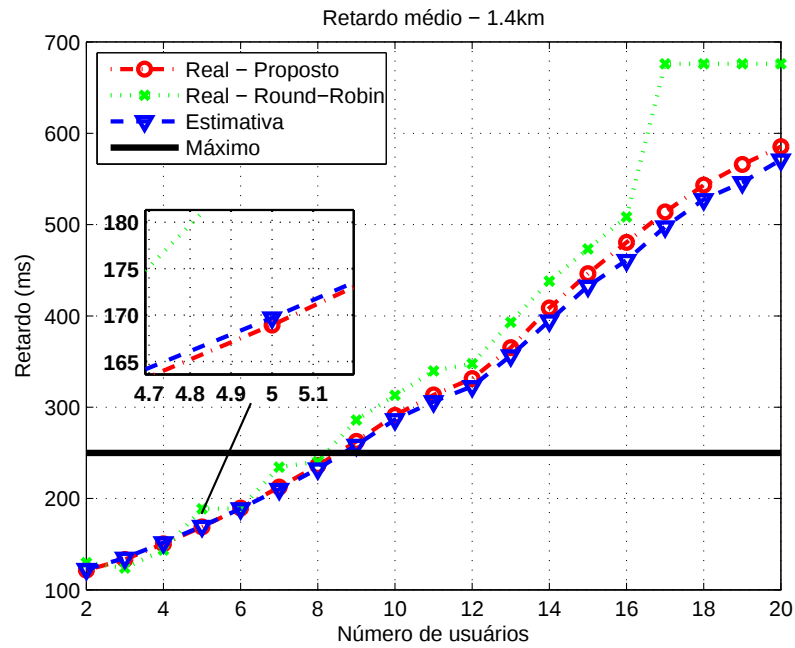


Figura 5.47: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1.4km, largura de banda de 5MHz, canal multipercurso Rayleigh e 1000 TTIs

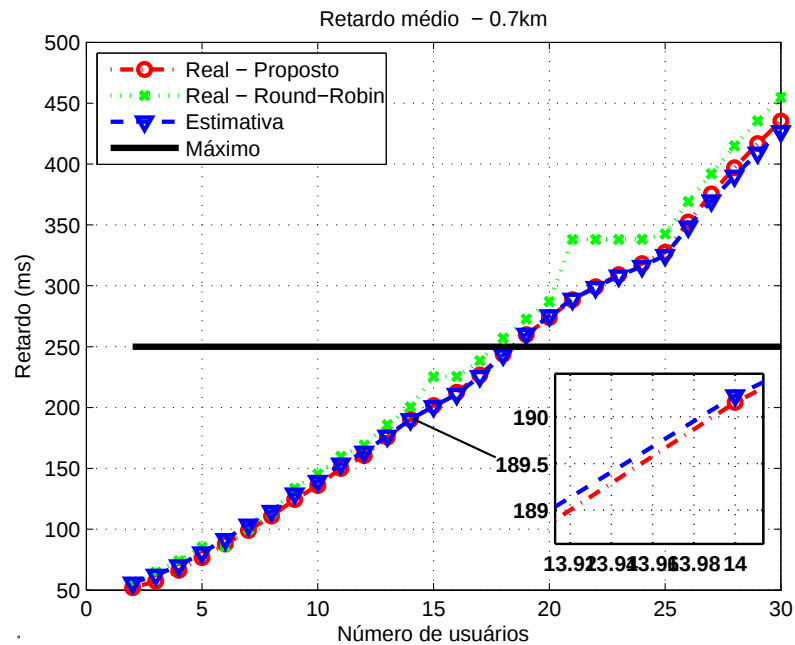


Figura 5.48: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 0.7km, largura de banda de 10MHz, canal multipercurso Rayleigh e 1000 TTIs

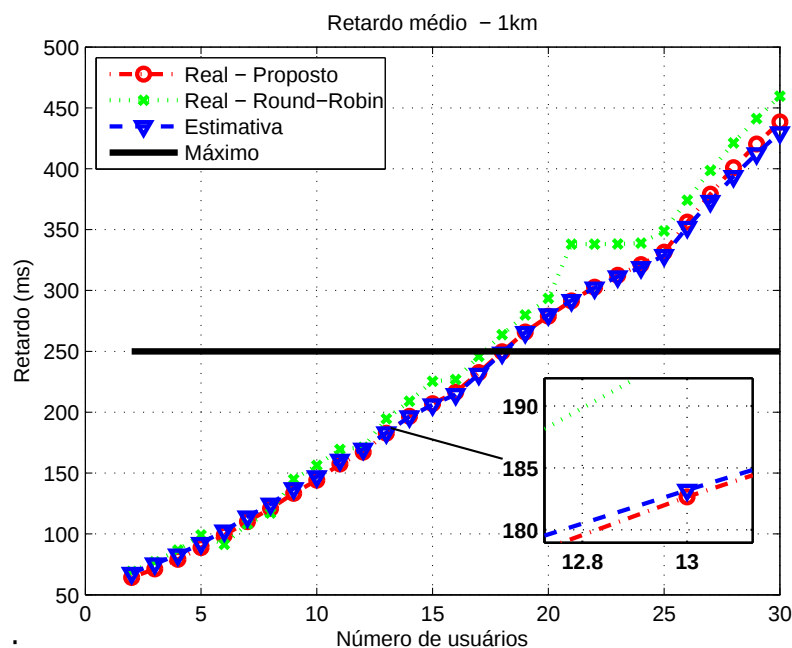


Figura 5.49: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

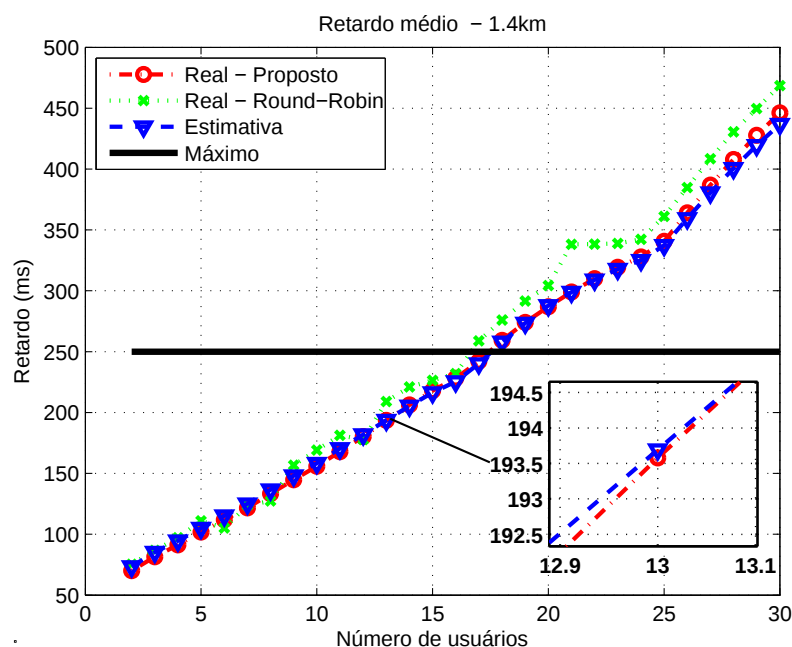


Figura 5.50: Retardo máximo tolerável, estimativa de limitante de retardo e retardo médio real calculado para algoritmo proposto e escalonador Round-Robin, considerando distância de 1.4km, largura de banda de 10MHz, canal multipercorso Rayleigh e 1000 TTIs

Conclusão

No Capítulo 4 é proposto um algoritmo que considera a qualidade de transmissão do canal e o critério de retardo máximo para decidir sobre a alocação de recursos de rádio disponíveis, buscando reduzir o retardo médio do sistema. A motivação de considerar o retardo máximo é que este é um parâmetro de QoS essencial para aplicações em tempo real com taxa de transmissão variável e requisitos específicos de banda como serviços de VoIP e de videoconferência, cada vez mais em evidência dado o crescimento exponencial do uso de dispositivos móveis.

Foram considerados cenários com diferentes modelagens de canal e largura de banda, a fim de testar o desempenho do algoritmo e validar a proposta comparando com o algoritmo PSO (*Particle Swarm Optimization*) [42] e o algoritmo com garantia de QoS [24]. O primeiro algoritmo objetiva maximizar a vazão total do sistema utilizando a heurística de otimização PSO, porém apresenta limitações em relação à complexidade computacional e ao retardo. Enquanto o algoritmo QoS garantido objetiva minimizar a complexidade computacional do problema de otimização da vazão total atendendo a restrição de taxa mínima de transmissão, porém com limitações em termos de taxa de perda. O algoritmo proposto é uma variação do algoritmo QoS garantido, porém com atenção ao critério de retardo ao invés do critério de taxa mínima de transmissão, e também com critérios de prioridade e de estimativa de blocos diferentes. Desta forma pretende-se garantir desempenho satisfatório em todos os quesitos de QoS estudados, tendo o algoritmo PSO e algoritmo QoS garantido como referência.

Verifica-se que o algoritmo proposto apresenta retardo médio do sistema inferior ao apresentado pelos algoritmos QoS garantido e PSO, além de garantir maior taxa de atendimento ao critério de retardo máximo definido, próximo de 100%. Observou-se que nos casos considerados o escalonamento (transmissão de dados) é realizado para todos os usuários em nossa proposta, ou seja, o retardo médio do sistema não tende ao infinito. Este fato se justifica pela característica de balanceamento do algoritmo proposto, ou seja, os usuários com melhores condições de canal têm maior número de blocos estimados para si, ao mesmo tempo que os usuários com piores condições de canal têm maior prioridade.

O algoritmo proposto apresenta desempenho em termos de vazão similar ao

algoritmo QoS garantido para número maior de usuários e inferior ao algoritmo PSO. O desempenho superior do algoritmo proposto para número maior de usuários se deve ao fato deste estimar número maior de blocos para usuários com melhores condições de canal, garantindo assim melhoria em termos de vazão. O fato do algoritmo PSO apresentar melhor desempenho no quesito vazão se justifica por se tratar de heurística de maximização da vazão em detrimento do tempo de processamento. Em relação ao tempo de processamento, o desempenho do algoritmo proposto é similar ao algoritmo QoS garantido.

Em relação ao índice de justiça o algoritmo proposto apresenta desempenho superior ao algoritmo PSO e inclusive superior ao algoritmo QoS garantido se considerado número elevado de usuários em simulação, em virtude de sua característica de balanceamento de usuários. Sua característica de balanceamento também permite ao algoritmo proposto apresentar taxa de perda média do sistema inferior ao algoritmo QoS garantido em todos os cenários e inferior ao algoritmo PSO para largura de banda de 10MHz.

No Capítulo 5 é proposta uma variação do algoritmo proposto no Capítulo 4, mantendo o objetivo de reduzir o retardo médio do sistema. Ao invés do retardo real da rede, é utilizada a informação de limitante de retardo estimado por Cálculo de Rede através de conceitos de curva de serviço e processo envelope MFBAP (*Multifractal Bounded Arrival Process*), este último escolhido por se tratar de uma heurística que se baseia em modelagem multifractal que apresenta resultados mais próximos ao envelope real. Assim, em nossa proposta, podemos atualizar a estimativa de retardo a medida que as características dos dados de tráfego no sistema variam a fim de tomar decisões antecipadas sobre a alocação de recursos na transmissão *downlink* LTE, principalmente em relação ao controle de admissão de usuários de maneira a atender requisitos de retardo tolerável.

Os resultados das simulações mostram que o desempenho do algoritmo proposto com estimativa de retardo é em geral superior ou similar ao desempenho dos algoritmos QoS garantido e PSO, além de melhorar o desempenho do algoritmo proposto, principalmente em termos de taxa de perda e retardo médio. O fato do algoritmo com estimativa utilizar valores limitantes de retardo possibilita a este melhoria nestes quesitos.

O algoritmo proposto com estimativa de retardo apresenta em geral menores valores para taxa de perda e retardo que os demais algoritmos simulados, além de garantir taxa de atendimento ao critério de retardo definido próximo de 100%, taxa superior aos demais algoritmos.

Em termos de vazão total e vazão média por usuário, o algoritmo proposto com estimativa de retardo provê desempenho similar ou superior ao algoritmo QoS garantido para número maior de usuários considerados em simulação e superior ao algoritmo proposto. Nota-se também que o algoritmo proposto com estimativa de retardo apresenta baixo tempo computacional e valores de *fairness* maiores do que os apresentados pelo

algoritmo PSO.

Em todos os cenários simulados, os valores de retardo estimados se mantêm próximos aos valores reais calculados. Dessa forma, antes mesmo de se iniciar o escalonamento dos recursos disponíveis, pode-se utilizar a informação do retardo estimado a fim de fornecer o número de usuários que podem ser admitidos para atender um certo valor máximo de retardo.

Como proposta de trabalho futuro, pretende-se obter uma curva de serviço específica para o escalonador proposto a fim de diminuir o erro quadrático médio da estimativa de limitante de retardo e verificar o desempenho em rede.

Referências Bibliográficas

- [1] 3GPP TS 36.104 version 8.3.0 Release 8. *LTE; Evolved Universal Terrestrial Radio Access (E-UTRA); Base Station (BS) radio transmission and reception*. s. 1., 2008.
- [2] 3GPP TR 25.913 version 8.0.0 Release 8. *Universal Mobile Telecommunications System (UMTS); LTE; Requirement for Evolved Universal Terrestrial Radio Access (E-UTRA) and Universal Terrestrial Radio Access Network (E-UTRAN)*, 2009.
- [3] 3GPP TR 36.931 version 9.0.0 Release 9. *LTE; Evolved Universal Terrestrial Radio Access (E-UTRA); Radio Frequency (RF) requirements for LTE Pico Node B*. s. 1., 2011.
- [4] Abrahão, D.C.; Vieira, F.H.T.; Ferreira, M.V.G., “Algoritmo de Alocação de Blocos de Recursos com Atendimento a Critério de Justiça para Redes sem Fio LTE”. In: XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014, Salvador, BA. Anais do XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014.
- [5] Abrahão, D.C.; Vieira, F.H.T.; Ferreira, M.V.G., “Resource Allocation Algorithm for LTE Networks Using β MWM Modeling and Adaptive Effective Bandwidth Estimation”. In: VI International Workshop on Telecommunications, 2015, Santa Rita do Sapucaí - MG. IEEE IWT, 2015.
- [6] Back, T.; Fogel, D. B.; Michalewicz, Z. *Handbook of Evolutionary Computation*. In: Smith, A. E.; Coit, D. W. *Constraint-Handling Techniques - Penalty Functions*. Pittsburgh: Taylor & Francis, 1997. p. C5.2:1- 1 C5.2:6.
- [7] Bae, Sueng Jae; Choi, Bum-Gon; Chung, Min Young, *Delay-Aware Packet Scheduling Algorithm for Multiple Traffic Classes in 3GPP LTE System*, Communications (APCC), 2011 17th Asia-Pacific Conference on , vol., no., pp.33,37, 2-5 Oct. 2011.
- [8] Bahai, A. R. S.; Saltzberg, B. R. *Multi-Carrier Digital Communications: Theory and Applications of OFDM*. 1. Ed. New York: Kluwer Academic, 2002.
- [9] Boudec, J.Y.L.; Thiran, P., *Network Calculus: A Theory of Deterministic Queuing Systems for the Internet*, Springer Verlag, Lecture Notes in Computer Science, LNCS 2050, 2004.

- [10] Cisco Systems, Inc.. Understanding Delay in Packet Voice Networks. Disponível em <<http://www.cisco.com/c/en/us/support/docs/voice/voice-quality/5125-delay-details.html>>, acesso em julho/2015.
- [11] Cho, Y. S. et al. *MIMO-OFDM Wireless Communications with MATLAB*. 1. Ed. Singapore: Wiley, 2010.
- [12] Chui, C. K.. *An Introduction to Wavelets*. San Diego: Academic, 1992.
- [13] Costa, V.H.T., *Análise de desempenho de sistemas de comunicação OFDM-TDMA utilizando cadeias de Markov e curva de serviço*, Goiânia - GO, 2013.
- [14] Cox, C. *An Introduction to LTE: LTE, LTE-Advanced, SAE, VoLTE and 4G Mobile Communications*. 2. Ed. Chinchester: Wiley, Ed. 2014.
- [15] Dahlman E.; Parkvall S.; Skold, J.; Beming P., *3G Evolution: HSPA and LTE for Mobile Broadband*, Elsevier, Great Britain, 2008.
- [16] Delgado, O.; Jaumard, B., "Scheduling and Resource Allocation in LTE Uplink with a Delay Requirement", Communication Networks and Services Research Conference (CNSR), 2010 Eighth Annual , vol., no., pp.268,275, 11-14 May 2010.
- [17] Duffield, N. G.; O Connell, N.. "Large deviations and overflow probabilities for the general single-server queue, with applications." Technical Report 1, Dublin Institute for Advanced Studies-Applied Probability Group, DIAS-STP-93-30, 1993.
- [18] Ferreira, M.V.G.; Vieira, F.H.T., "Algoritmo de Alocação de Blocos de Recursos em Redes LTE com Garantia de Retardo Máximo aos Usuários". 11º Congresso de Pesquisa, Ensino e Extensão (Conpeex). Goiânia. 2014.
- [19] Ferreira, M.V.G.; Vieira, F.H.T; Abrahão, D.C., "Minimizing Delay in Resource Block Allocation Algorithm of LTE Downlink", VI International Workshop on Telecommunications: IWT 2015.
- [20] Ferreira, M.V.G.; Vieira, F.H.T.; Castro, M.; Araújo, S.; Rocha, F.G.C., "Alocação de Blocos de Recurso em Redes LTE com Estimativa de Limitante de Retardo através de Cálculo de Rede". In: XXXIII SBRT 2015, Juiz de Fora - MG. Anais do XXXIII SBRT, 2015.
- [21] Fette, B. et al. *RF & Wireless Technologies: Know It All*. 1 Ed. Chennai: Elsevier, 2007.
- [22] FracLab 2.0. General purpose signal and image processing toolbox based on fractal and multifractal methods. Disponível em <<http://fraclab.saclay.inria.fr/>>, acesso em agosto/2015.

- [23] Garg, V. K. *Wireless Communications and Networking*. 1. Ed. San Francisco: Elsevier, 2007.
- [24] Guan, Na; Zhou, Yiqing; Tian, Lin; Sun, Gang; Shi, Jinglin, “QoS guaranteed resource block allocation algorithm for LTE systems”, WiMob, 2011 IEEE 7th International Conference on , vol., no., pp.307,312, 10-12 Oct. 2011.
- [25] Holma, H; Toskala, A., *LTE for UMTS - OFDMA and SC-FDMA Based Radio Access*, John Wiley & Sons, Ltd. ISBN: 978-0-470-99401-6, 2009.
- [26] Huang, Jeng-Ji; Lin, Wei-Keng; Ko, Hung-Hsiang, “A Resource Allocation Algorithm for Maximizing Packet Transmission in Downlink LTE Cellular Systems”, TENCON 2011 - 2011 IEEE Region 10 Conference , vol., no., pp.445,449, 21-24 Nov. 2011.
- [27] International Telecommunication Union (ITU). G.114 : One-way transmission time. Disponível em <<https://www.itu.int/rec/T-REC-G.114/en>>, acesso em julho/2015.
- [28] Jain, Raj; Durrezi, Arjan; Babic, Gojko, *Throughput Fairness Index: An Explanation*, Department of CIS, The Ohio State University, ATM_Forum/99-0045, 1999.
- [29] Kelly, F.. *Notes on effective bandwidths. In Stochastic Networks*, Oxford University Press, 1996.
- [30] Kennedy, J.; Eberhart, R. “Particle Swarm Optimization”. Proceedings of IEEE International Conference on Neural Networks IV. Washington DC. p. 1942–1948. 1995.
- [31] Lima, A. B. *Contribuições à modelagem de teletráfego fractal*. São Paulo, 2008.
- [32] Melo, C. A. V.; Fonseca, N. L. S. “An envelope process for multifractal traffic modeling”. v. 4, pp. 2168-2173. Communications, 2004 IEEE International Conference on. June, 2004.
- [33] Ni, Mingjian; Xu, Xiang; Mathar, R., “A Channel Feedback Model with Robust SINR Prediction for LTE Systems”, 7th European Conference on Antennas and Propagation (EuCAP), 2013.
- [34] Pereira, F. M., *Policiamento e escalonamento de tráfego em redes Ethernet PON*, Campinas - SP, 2006.
- [35] Prasad, R. *OFDM for Wireless Communications Systems*. 1. Ed. London: Artech House, 2004.

- [36] Ribeiro, V. J.; Riedi, R. H.; Crouse, M. S.; Baraniuk, R. G.. “Multiscale queueing analysis of long-range dependent traffic”. Proc. IEEE Infocom, 2000.
- [37] Riedi, R. H.; Crouse, M. S.; Ribeiro, V. J. e Baraniuk, R. G.. “A multifractal wavelet model with application to network traffic.” IEEE Trans. on Information Theory, vol. 45, no.3, pp. 992-1018, 1999.
- [38] Rocha, F. G. C.; Vieira, F. H. T.. *Modelagem de tráfego de vídeo MPEG-4 utilizando cascata multifractal com distribuição autorregressiva dos multiplicadores*. I2TS, Florianópolis, SC, 2009.
- [39] Sauter, M. *From GSM to LTE: An introduction to mobile networks and mobile broadband*. 1. Ed. Chichester: Wiley, 2011.
- [40] Sousa, P.B.. *Comparação de Esquemas de Alocação de Recursos para Downlink de Redes LTE*. Goiânia - Goiás, 2014.
- [41] Spiegel, M. R.; Liu, J.. *Manual de Fórmulas e Tabelas Matemáticas*. Col. Schaum - 3ª Ed, Makron Books, 2011.
- [42] Su, Lin; Wang, Ping; Liu, Fuqiang, “Particle swarm optimization based resource block allocation algorithm for downlink LTE systems”, Communications (APCC), 2012 18th Asia-Pacific Conference on , vol., no., pp.970,974, 15-17 Oct. 2012.
- [43] Tarokh, V. *New Directions in Wireless Communications Research*. 1. Ed. Cambridge: Springer, 2009.
- [44] Vieira, F.H.T.; Santos, J.A.; Cardoso, A.A., “Estimation of backlog and delay in OFDM/TDMA systems with traffic policing using Network Calculus”, Computers and Electrical Engineering. Volume 39, Issue 8, pages 2507–2520, November 2013.
- [45] Vieira, F.H.T.; Gonçalves, B.H.P.; Ferreira, M.V.G., “Algoritmo Baseado em Enxame de Partículas e Modelagem de Tráfego β MWM para Alocação Dinâmica de Recursos em Redes LTE”. In: XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014, Salvador, BA. Anais do XLVI Simpósio Brasileiro de Pesquisa Operacional, 2014.
- [46] Vieira, F.H.T.; Gonçalves, B.H.P.; Rocha, F.G.C.; Lee, L.L.; Ferreira, M.V.G., “Dynamic Resource Allocation in LTE Systems using an Algorithm based on Particle Swarm Optimization and BetaMWM Network Traffic Modeling”. 6th IEEE Latin American Symposium on Circuits and Systems: IEEE LASCAS 2015.
- [47] Yahiya, T. A.; *Understanding LTE and its Performance*. 1. Ed. Orsay Cedex: Springer, 2011.

- [48] Wand Network Research Group, University of Waikato. Séries reais de tráfego TCP/IP. Disponível em <<http://wand.net.nz/wits/waikato/8/>>, acesso em julho/2015.
- [49] Wang, J.H.; Yin, Z.Y.. “A ranking selection-based particle swarm optimizer for engineering design optimization problems”. Structural and Multidisciplinary Optimization, Volume 37, Issue 2, pp 131-147, 2008.

Séries Reais de Tráfego TCP/IP - Universidade de Waikato

Foram consideradas cinco séries reais de tráfego TCP/IP (*Transmission Control Protocol/Internet Protocol*) para representação dos usuários durante a simulação, agregadas no domínio do tempo em intervalos de 100ms, conforme ilustrado nas Figuras A.1, A.2, A.3, A.4 e A.5.

As Tabelas A.1, A.2, A.3, A.4 e A.5 apresentam os dados estatísticos referentes às series de tráfego. As séries reais consideradas representam o tráfego TCP/IP entre a Universidade de Waikato com redes externas, coletados entre 20/05/2011 e 29/10/2011 [48].

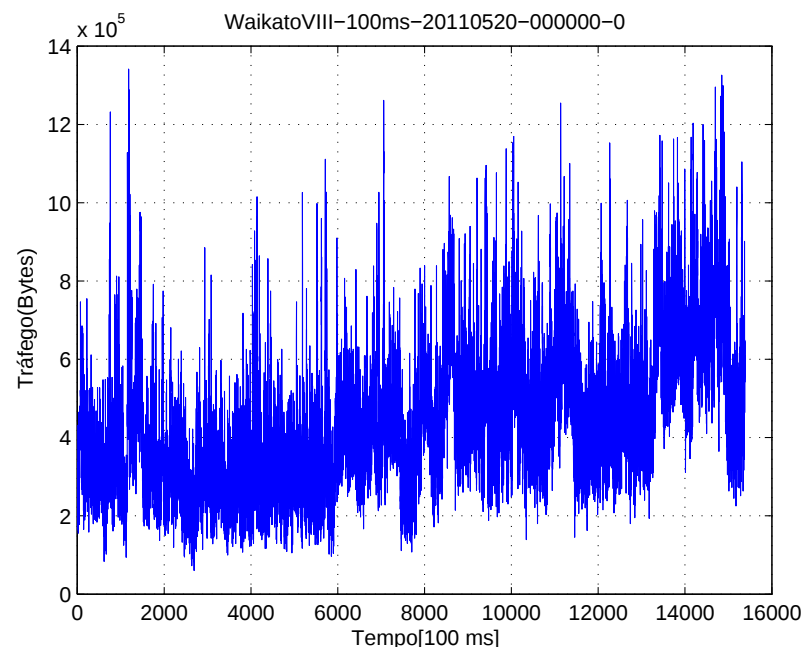


Figura A.1: Série real de tráfego TCP/IP WaikatoVIII-100ms-20110520-000000-0 [48] considerada para usuários 1, 5, 10, 15, 20, 25 e 30

Tabela A.1: *Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110520-000000-0 [48]*

Estatística	Valor
Média	438934
Desvio padrão	181968
Variância	3.31124e+10
Valor mínimo	60716
Valor máximo	1.34137e+06
Índice de dispersão	75438.2
Variância normalizada	0.171867
Momento de 3ª ordem (simetria)	5.67997e+15
Momento de 4ª ordem (concentração)	4.41719e+21

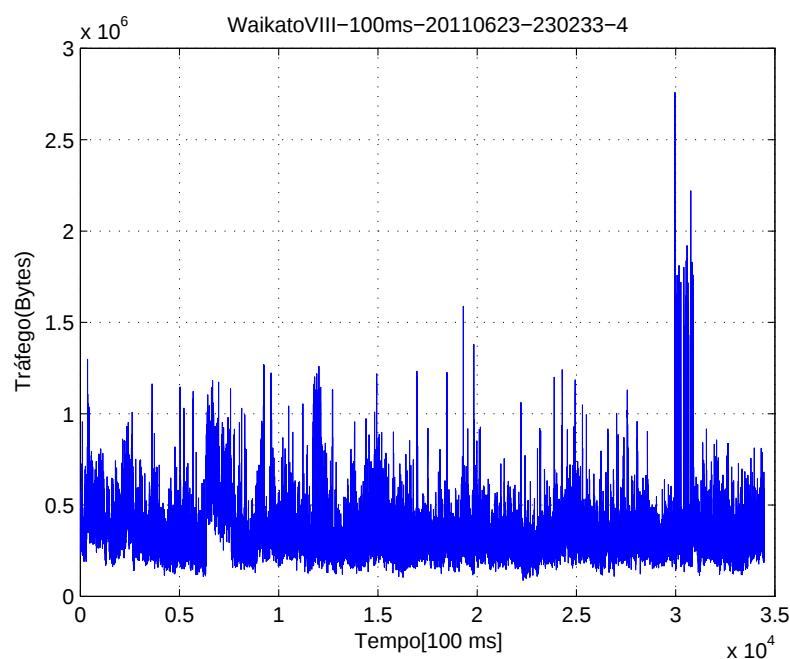


Figura A.2: *Série real de tráfego TCP/IP WaikatoVIII-100ms-20110623-230233-4 [48] considerada para usuários 2, 6, 11, 16, 21 e 26*

Tabela A.2: *Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110623-230233-4 [48]*

Estatística	Valor
Média	348920
Desvio padrão	162313
Variância	2.63456e+10
Valor mínimo	88542
Valor máximo	2.75798e+06
Índice de dispersão	75506.1
Variância normalizada	0.216399
Momento de 3ª ordem (simetria)	1.29634e+16
Momento de 4ª ordem (concentração)	1.34117e+22

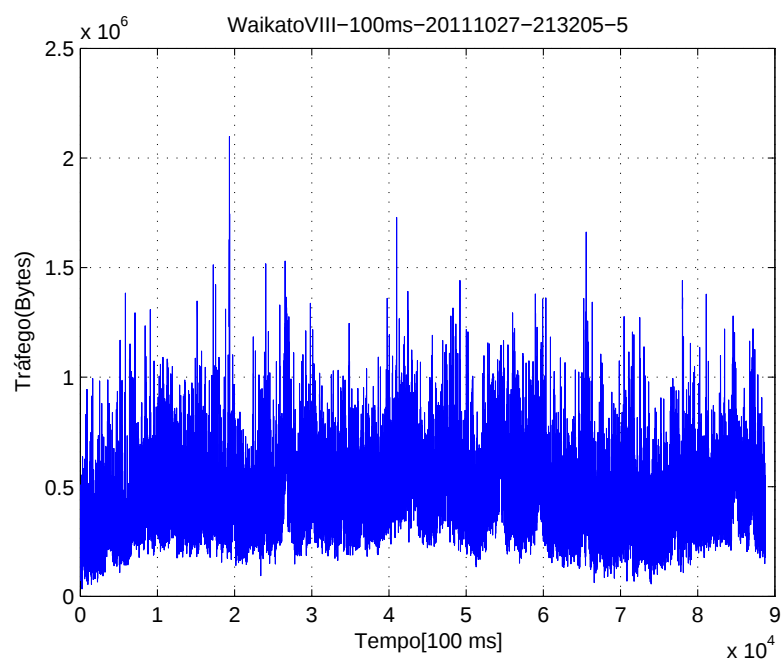


Figura A.3: *Série real de tráfego TCP/IP WaikatoVIII-100ms-2011027-213205-5 [48] considerada para usuários 3, 7, 12, 17, 22 e 27*

Tabela A.3: *Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20111027-213205-5 [48]*

Estatística	Valor
Média	459749
Desvio padrão	164831
Variância	2.71694e+10
Valor mínimo	34505
Valor máximo	2.09772e+06
Índice de dispersão	59096.2
Variância normalizada	0.12854
Momento de 3ª ordem (simetria)	4.4362e+15
Momento de 4ª ordem (concentração)	3.846e+21

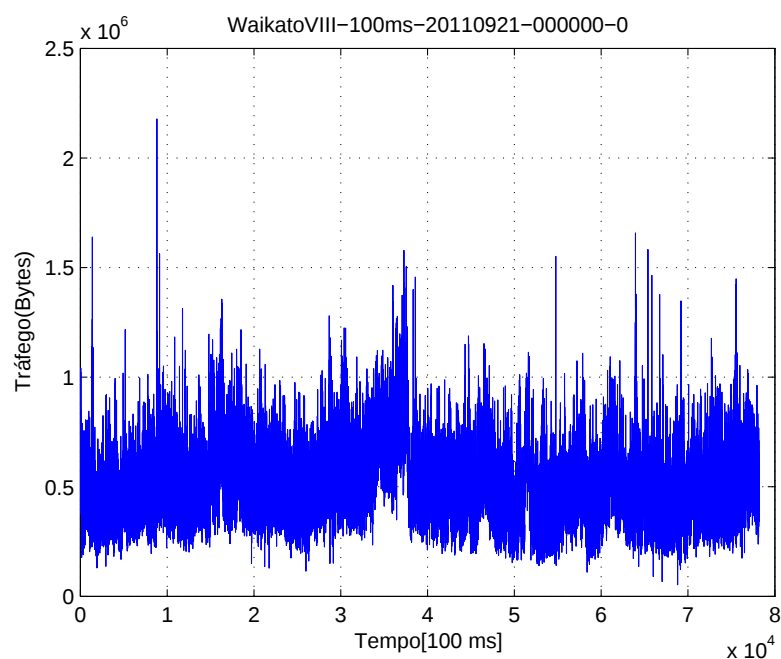


Figura A.4: *Série real de tráfego TCP/IP WaikatoVIII-100ms-20110921-000000-0 [48] considerada para usuários 4, 8, 13, 18, 23 e 28*

Tabela A.4: *Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20110921-000000-0 [48]*

Estatística	Valor
Média	497209
Desvio padrão	160166
Variância	2.56531e+10
Valor mínimo	52773
Valor máximo	2.17751e+06
Índice de dispersão	51594.1
Variância normalizada	0.103768
Momento de 3ª ordem (simetria)	4.12028e+15
Momento de 4ª ordem (concentração)	3.11196e+21

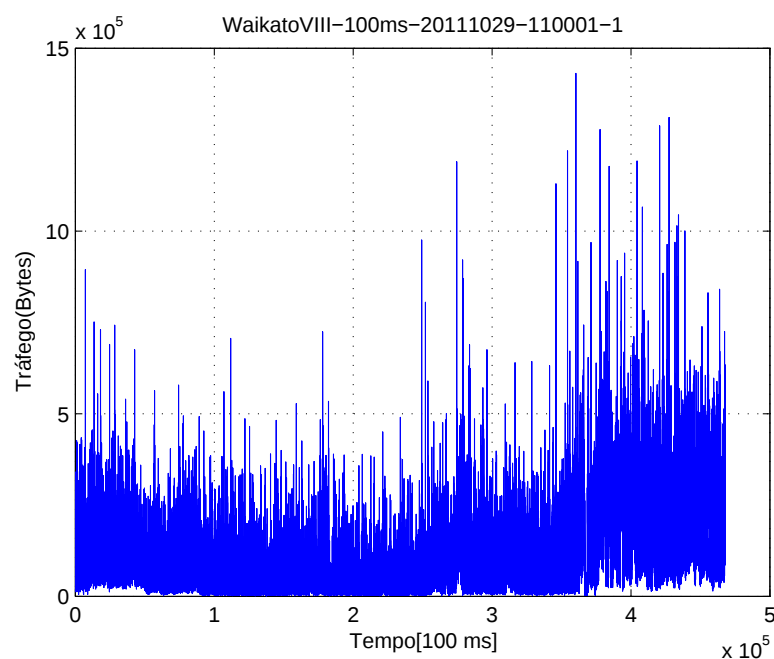


Figura A.5: *Série real de tráfego TCP/IP WaikatoVIII-100ms-2011029-110001-1 [48] considerada para usuários 5, 9, 14, 19, 24 e 29*

Tabela A.5: *Dados estatísticos referentes à série real de tráfego TCP/IP WaikatoVIII-100ms-20111029-110001-1 [48]*

Estatística	Valor
Média	100934
Desvio padrão	87453.3
Variância	7.64807e+09
Valor mínimo	610
Valor máximo	1.43172e+06
Índice de dispersão	75773.2
Variância normalizada	0.750722
Momento de 3ª ordem (simetria)	1.35246e+15
Momento de 4ª ordem (concentração)	7.17766e+20