

# ENERGY-AWARE RELIABLE COMMUNICATION

by

HAJAR MAHDAVI-DOOST

A dissertation submitted to the

Graduate School—New Brunswick

Rutgers, The State University of New Jersey

In partial fulfillment of the requirements

For the degree of

Doctor of Philosophy

Graduate Program in Electrical and Computer Engineering

Written under the direction of

Roy D. Yates

And approved by

---

---

---

---

---

New Brunswick, New Jersey

May, 2016

© 2016

Hajar Mahdavi-Doost

**ALL RIGHTS RESERVED**

## **ABSTRACT OF THE DISSERTATION**

### **Energy-Aware Reliable Communication**

**By HAJAR MAHDAVI-DOOST**

**Dissertation Director:**

**Roy D. Yates**

Emerging applications of short-range communication such as the Internet of Things and body area networks highlight the importance of processing energy, as compared to transmit energy. In this thesis, we investigate fundamental limits of reliable communication when receiver processing is powered by random energy sources and subject to constraints on energy storage. We propose a receiver model that captures the trade-off between sampling energy and decoding energy. The model relies on the decoding energy being a decreasing function of the capacity gap between the code rate and the channel capacity. The receiver can save energy in sampling by dropping a fraction of samples, at the cost of reducing the effective capacity and thus increasing the energy needed for decoding. While sampling and decoding energies are typically comparable, the key issue is that the sampling is a real-time process; the samples must be collected during the transmission time of that packet. Thus the energy harvesting rate and battery size may constrain the sampling rate. This model allows us to characterize the maximum throughput of a basic communication channel with limited processing energy. This is done based on striking the balance between the sampling and decoding energy, subject

to limited random arrival of energy, and limited battery size.

We further extend this result to multi-user scenarios, where multiple transmitters communicate with a single receiver with limited energy. We introduce the concept of receive multi-user diversity, in which the receiver decodes the messages experiencing the strongest channels in order to reduce the decoding energy per user.

Next, we propose using hybrid automatic retransmission request (HARQ) with soft combining to reduce the processing energy and improve the throughput under limited receiver energy. In this protocol, the receiver keeps requesting additional redundancy in order to increase the capacity gap, which in turn reduces the processing energy. We compare the performance of incremental redundancy (IR) HARQ, and Repetition-HARQ. In these systems, the decoding energy is a decreasing function of the capacity gap but an increasing function of the code-length. The IR-HARQ protocol yields a better capacity gap, but increases the code-length, while Repetition-HARQ offers less improvement in the capacity gap, but does not increase the effective code-length. Thus, contrary to systems without receiver energy constraints in which IR-HARQ always performs better, here, depending on the system parameters, Repetition-HARQ can outperform IR-HARQ.

Finally, we study energy efficiency and energy harvesting in LTE networks. We formulate a single-cell downlink scheduling problem that enforces constraints on the selection of transmission parameters. Linear cost constraints on the set of channels are also imposed in order to accommodate energy efficiency considerations. We show that the resulting problem is NP-hard and we propose a deterministic multiplicative-update algorithm for which we establish an approximation guarantee. We also consider the problem of downlink scheduling in an LTE network powered by energy harvesting devices. We formulate optimization problems that seek to maximize two popular energy efficiency metrics subject to LTE network constraints and energy harvesting causality constraints. We focus on a key sub-problem and show that this problem is NP-hard. Then we reformulate it as a constrained submodular set function maximization problem which can be solved with a constant-factor approximation using a greedy algorithm.

## Acknowledgements

First of all, I would like to thank my advisor, Professor Roy D. Yates, for his invaluable guidance and support, to complete my thesis. In particular, he patiently taught me to precisely and rigorously articulate my mathematical arguments. During my years as a graduate student, he gave me space to grow into an independent researcher, aiming high, and not giving up. One of the most important things that I learned from Roy is that every challenging obstacle is indeed an opportunity for an interesting research initiative.

I would like to express my sincere gratitude to Dr. Narayan Prasad. I was fortunate to have him as my mentor during the course of my internship in NEC Labs America. He also continued working with me afterwards supervising the work done as part of my thesis. During the last year, in a close and daily-basis collaboration, he shared with me his deep and diverse knowledge on algorithm design, discrete optimization, and wireless standards. He has been a great mentor in my research and a role model for hard work and dedication.

I also would like to thank Professor Dipankar Raychaudhuri, Rich Howard and Rich Martin for their thoughtful comments and discussions throughout my research and also Ivan Seskar for his technical support and help with the resources in WINLAB. I would like to thank Professor Athina Petropulu and Emina Soljanin for their guidance, support, and encouragement. They are not only extraordinary researchers but also great inspirational female role models. I am also grateful for the comments I received from Professor Narayan Mandayam and Predrag Spasojevic on my research. I wish to thank my colleagues in WINLAB and ECE department at Rutgers who created a friendly and comfortable working environment. I gained great insight into my research through several discussions I had with many of them.

I would like to give special thank to my NJ friends who have been there always for me. They made the years I spent in Ph.D. program an amazing and unforgettable experience. I couldn't complete this work without their help and support.

The greatest acknowledgment goes to my family, my parents and my brothers, for their tremendous support and encouragement. In particular, I would thank my mom and dad for their unconditional love and belief they had in me, from my early childhood. They instilled in me a desire to learn and made sacrifices so I would have access to high quality education to follow my dreams.

Last but not least, I would like to extend my deepest gratitude, to the love of my life, Mohammad, for his tremendous love, patience, support, and understanding during the course of my study. His encouragement and optimism gave me motivation and hope through the ups and downs of my Ph.D. journey. Thank you Mohammad for standing by my side in this journey. Without your sacrifices and supports, I wouldn't be able to complete this thesis. I also would like to thank my wonderful daughter, Saba, for being such a good girl and always cheering me up. Her understanding and patience during my graduate study was beyond my expectations. Her endless joy and passion always gave me energy to continue my work and finally complete my thesis.

# Table of Contents

|                                                                          |    |
|--------------------------------------------------------------------------|----|
| <b>Abstract</b> . . . . .                                                | ii |
| <b>Acknowledgements</b> . . . . .                                        | iv |
| <b>List of Figures</b> . . . . .                                         | ix |
| <b>1. INTRODUCTION</b> . . . . .                                         | 1  |
| 1.1. Overview . . . . .                                                  | 1  |
| 1.2. Energy Trade-off at the Receiver . . . . .                          | 3  |
| 1.3. Receiver Sampling and Decoding: Power Comparisons . . . . .         | 7  |
| 1.4. Notations . . . . .                                                 | 8  |
| 1.5. Thesis Outline . . . . .                                            | 8  |
| <b>2. Energy Harvesting Receivers: Finite Battery Capacity</b> . . . . . | 11 |
| 2.1. System Model . . . . .                                              | 11 |
| 2.1.1. Energy Harvesting Model . . . . .                                 | 12 |
| 2.1.2. Variable-Timing Transmission . . . . .                            | 13 |
| 2.1.3. Performance Metrics . . . . .                                     | 14 |
| 2.1.4. Preliminaries: Chernoff Hoeffding Inequality . . . . .            | 15 |
| 2.2. Achievability: Deterministic Energy Arrivals . . . . .              | 16 |
| 2.2.1. Low Harvesting Rate . . . . .                                     | 16 |
| 2.2.2. High Harvesting Rate . . . . .                                    | 18 |
| 2.3. Achievability: Random Energy Arrivals . . . . .                     | 20 |
| 2.3.1. Low Harvesting Rate . . . . .                                     | 21 |
| 2.3.2. High Harvesting Rate . . . . .                                    | 36 |
| 2.4. Outerbound . . . . .                                                | 38 |

|                                                                                                                       |           |
|-----------------------------------------------------------------------------------------------------------------------|-----------|
| 2.5. Optimum Code Rate . . . . .                                                                                      | 43        |
| <b>3. Energy-Harvesting Receivers in Fading Channels . . . . .</b>                                                    | <b>47</b> |
| 3.1. System Model . . . . .                                                                                           | 47        |
| 3.1.1. Energy Harvesting Model . . . . .                                                                              | 49        |
| 3.2. Achievability: Channel Selective Sampling . . . . .                                                              | 49        |
| 3.3. Variable Timing: Achievable Rates . . . . .                                                                      | 53        |
| 3.4. Outerbound . . . . .                                                                                             | 55        |
| 3.5. Conclusion . . . . .                                                                                             | 59        |
| <b>4. Opportunistic Reception in a Multiuser Slow-Fading Channel with an<br/>Energy Harvesting Receiver . . . . .</b> | <b>60</b> |
| 4.1. System Model . . . . .                                                                                           | 60        |
| 4.2. Single User System . . . . .                                                                                     | 62        |
| 4.2.1. Throughput Optimization . . . . .                                                                              | 63        |
| 4.2.2. Single User Performance/ Self-imposed Threshold . . . . .                                                      | 69        |
| 4.3. Multiuser System . . . . .                                                                                       | 71        |
| 4.4. Conclusions . . . . .                                                                                            | 72        |
| <b>5. Hybrid ARQ in Block-Fading Channels with an Energy Harvesting<br/>Receiver . . . . .</b>                        | <b>74</b> |
| 5.1. System Model . . . . .                                                                                           | 74        |
| 5.2. ARQ Schemes . . . . .                                                                                            | 76        |
| Classic ARQ . . . . .                                                                                                 | 76        |
| Repetition-HARQ . . . . .                                                                                             | 77        |
| IR-HARQ . . . . .                                                                                                     | 77        |
| 5.2.1. Objective Function . . . . .                                                                                   | 77        |
| 5.3. Decoding Energy . . . . .                                                                                        | 78        |
| Classic ARQ . . . . .                                                                                                 | 78        |
| Repetition-HARQ . . . . .                                                                                             | 78        |



|                                                                                                   |            |
|---------------------------------------------------------------------------------------------------|------------|
| IR-HARQ . . . . .                                                                                 | 79         |
| IR-HARQ with subset selection . . . . .                                                           | 80         |
| 5.4. Decision Policy . . . . .                                                                    | 81         |
| 5.5. Simulation Results . . . . .                                                                 | 82         |
| <b>6. Energy-Aware Downlink Scheduling in LTE-Advanced Networks . .</b>                           | <b>86</b>  |
| 6.1. Problem formulation . . . . .                                                                | 88         |
| 6.1.1. Optimization framework . . . . .                                                           | 88         |
| 6.1.2. Backlogged traffic model . . . . .                                                         | 91         |
| 6.1.3. Finite queue model . . . . .                                                               | 92         |
| 6.1.4. Hardness result . . . . .                                                                  | 92         |
| 6.2. A Unified Scheduling Algorithm . . . . .                                                     | 93         |
| 6.2.1. Benchmarking . . . . .                                                                     | 108        |
| 6.3. Simulation Results . . . . .                                                                 | 110        |
| 6.4. Conclusions . . . . .                                                                        | 113        |
| <b>Appendices . . . . .</b>                                                                       | <b>114</b> |
| <b>7. Optimizing Energy Efficiency over Energy-Harvesting LTE Cellular<br/>Networks . . . . .</b> | <b>121</b> |
| 7.1. Problem formulation . . . . .                                                                | 122        |
| 7.1.1. Practical Constraints . . . . .                                                            | 123        |
| 7.1.2. Objective Functions . . . . .                                                              | 126        |
| 7.2. A Constant Factor Approximation Algorithm . . . . .                                          | 130        |
| 7.3. Extensions . . . . .                                                                         | 139        |
| 7.4. Simulation Results . . . . .                                                                 | 143        |
| 7.5. Conclusions . . . . .                                                                        | 148        |
| <b>8. Conclusions and Future Work . . . . .</b>                                                   | <b>149</b> |
| <b>References . . . . .</b>                                                                       | <b>152</b> |

## List of Figures

|      |                                                                                                                                                                                                                                                                                                                                                                                                                                                                                                  |    |
|------|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 1.1. | Number of iterations versus the sampling rate in the belief propagation iterative algorithm in LDPC code from DVB-S.2 with block length $n = 64,800$ and code rate $R = 3/4$ bits/s/Hz. . . . .                                                                                                                                                                                                                                                                                                  | 4  |
| 1.2. | Number of iterations versus the sampling rate in the belief propagation iterative algorithm in an Irregular LDPC decoder based on Richardson and Urbanke's 2001 paper, $n = 1000$ , $R = 1/2$ bits/sec/Hz. . . . .                                                                                                                                                                                                                                                                               | 5  |
| 1.3. | Normalized decoding energy $f(\lambda) = \mathcal{E}_D(\lambda)$ for a fixed $R$ and $C$ as function of the sampling rate $\lambda$ . The total energy per symbol of the receiver is minimized at $\lambda = \lambda^*$ . . . . .                                                                                                                                                                                                                                                                | 6  |
| 2.1. | Stored energy (in $\mu J$ ) in different time intervals of variable-timing transmission with random energy harvesting. The intervals marked "C," "S," and "D" mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. Here $W_t = 5 \mu J$ or $W_t = 1 \mu J$ with probability 0.1 each and $W_t = 0 \mu J$ with probability 0.8 while $\bar{W} = 0.6 \mu J$ . Also, $\beta = 0.3 \mu J$ , $n = 1000$ and the sampling rate $\lambda = 0.9$ . . . . . | 13 |
| 2.2. | Variable-timing optimum policy: Transmitted packets are labeled $\mathbf{x}_1, \mathbf{x}_2, \dots$ while intervals marked "C", "S", and "D" mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. The corresponding graph depicts the receiver's stored energy. . . . .                                                                                                                                                                            | 16 |
| 2.3. | Stored energy over the sampling interval modeled as the energy queue. Comparing to Fig. 2.1 and Fig. 2.2, here just a single (S)ampling interval is shown. . . . .                                                                                                                                                                                                                                                                                                                               | 23 |
| 2.4. | Total energy as a function of the sampling rate when $\hat{\lambda}_n = \beta + \bar{W} + n^{\alpha-1} < \lambda^*$ . . . . .                                                                                                                                                                                                                                                                                                                                                                    | 41 |

|                                                                                                                                                                                                                                                                                                                    |    |
|--------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 2.5. Sampling rate, communicate rate and the total normalized energy versus $R$ (bits/s/Hz) for $\bar{W} = 0.3\mu J$ , $\beta = 0.5\mu J$ , $C = 10$ bits/s/Hz, $n = 10^6$ .                                                                                                                                       | 46 |
| 3.1. Channel Model                                                                                                                                                                                                                                                                                                 | 48 |
| 3.2. Variable-timing optimum policy: Transmitted packets are labeled $\mathbf{x}_1, \mathbf{x}_2, \dots$ while intervals marked “C,” “S,” and “D” mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. The corresponding graph depicts the receiver’s stored energy. | 49 |
| 3.3. Capacity of the effective channel under channel selective sampling versus the threshold gain in Rayleigh fading.                                                                                                                                                                                              | 52 |
| 3.4. Normalized decoding energy, $\mathcal{E}_D(\gamma)$ , versus sampling energy, $\mathcal{E}_S(\gamma) = \Pr[G \geq \gamma]$ for a fixed $R$ under channel selective sampling strategy. The total energy requirement is minimum at Minimum Energy Point.                                                        | 53 |
| 3.5. The achievable rate for channel selective sampling strategy and the strategy at which the sampling is optimized independently from channel with optimum parameters in terms of the transmit power.                                                                                                            | 55 |
| 4.1. Optimum threshold $\tilde{\gamma}$ and optimum throughput versus $n\bar{W}$ for $R = 0.3$ bits/s/Hz, $\nu = 1.69$ , $\eta = 0.42$ .                                                                                                                                                                           | 69 |
| 4.2. Comparing the achievable service rate of the opportunistic selection with what suggested in (4.34), $n = 10^6$ .                                                                                                                                                                                              | 70 |
| 4.3. Comparison of the number of users being served in the opportunistic and random selection for two $\bar{W} = 100 \mu J$ and $\bar{W} = 1000 \mu J$ when $R = 0.3$ bits/s/Hz, $\nu = 1.69$ , $\eta = 0.42$ , $n = 10^6$ .                                                                                       | 71 |
| 4.4. The sum-rate versus the number of users for $R = 0.3, 0.6, 0.9$ bits/s/Hz when $\eta = 0.1$ , $\nu = 2.5 \times 10^{-4}$ and $\bar{W} = 100 \mu J$ , $n = 10^8$ .                                                                                                                                             | 73 |
| 5.1. Throughput versus the arrival energy rate of $n\bar{W}$ for the code rate of $R = 3$ bits/s/Hz, $n = 10^6$ .                                                                                                                                                                                                  | 83 |
| 5.2. Throughput versus the arrival energy rate of $n\bar{W}$ for the code rate of $R = 0.5$ bits/s/Hz, $n = 10^6$ .                                                                                                                                                                                                | 84 |

|                                                                                                                                                                                   |     |
|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 5.3. Throughput versus the code rate $R$ for the energy arrival $n\bar{W} = 0.3$<br>J/slot. . . . .                                                                               | 84  |
| 5.4. Throughput versus the code rate $R$ for the energy arrival $n\bar{W} = 2$ J/slot.                                                                                            | 85  |
| 6.1. Weighted sum rate vs number of RBs for $\lambda = 10, \theta = 10$ . . . . .                                                                                                 | 112 |
| 6.2. Weighted sum rate vs number of users for $\lambda = 10, \theta = 10$ . . . . .                                                                                               | 112 |
| 6.3. Comparison of weighted sum rate vs number of RBs for the case with<br>$\lambda = 10$ and $\theta = 10$ , and the case with the optimized $\lambda$ and $\theta = 10$ . . . . | 112 |
| 7.1. A feasible allocation for a system with $N = 4$ RBs, $L = 3$ subframes per<br>block, $K = 2$ users and $M = 3$ modes. . . . .                                                | 124 |
| 7.1. Achieved GEE versus Number of used RBs: Large Queue-sizes . . . . .                                                                                                          | 142 |
| 7.2. Achieved GEE versus Number of used RBs: Small Queue-sizes . . . . .                                                                                                          | 143 |
| 7.3. Achieved GEE versus Number of used RBs: Large Queue-sizes . . . . .                                                                                                          | 144 |
| 7.4. Achieved GEE versus Number of used RBs: Large Queue-sizes and $L = 3$<br>subframes per block . . . . .                                                                       | 145 |
| 7.5. Achieved GEE versus Number of used RBs: Small Queue-sizes and $L = 3$<br>subframes per block . . . . .                                                                       | 146 |

## Chapter 1

# INTRODUCTION

### 1.1 Overview

Energy harvesting offers the promise of unbounded lifetime extension to battery powered devices; however, the randomness of the energy source and it being limited have introduced new challenges. Recently, there has been considerable research on communication systems that rely on energy harvesting at the transmitter; see, for example, [1–6] for point-to-point channels and [7–13] for small networks. A survey of problem formulations appears in [14]. Although circuit and processing costs have been addressed in [15] and [16], the primary emphasis has been on the energy costs of transmission. Even in networks with energy harvesting relays [10, 12, 13], the energy cost of receiving at the relays is ignored. However, communication over short distances can achieve high rates with relatively small transmit power. In this case, the energy consumption associated with the complex detection and decoding operations of the receiver becomes the dominant system constraint. Yet, with some exceptions [17], there has been little work on the problem of energy harvesting at the receiver. Known results are technology dependent; for example, practical energy consumption models of the receiver front-end have appeared in [18, 19]. This thesis mainly addresses the problem of energy harvesting at the receiver which is based on our publications [20–24]. Shortly after our work, some papers appeared on optimizing the communication scheme when both transmitter and receiver are harvesting energy. They formulate an optimization problem and solve it using the convexity of the decoding energy in terms of the code rate [25, 26]. However, our model is beyond the model assumed in these papers in the sense that we will also consider the sampling energy expenditure. In terms of the decoding energy, we will

compare our model with the practical decoding schemes like LDPC codes. We also mainly optimize the communication under a fixed code rate assumption although we also discuss optimizing of the code rate as well.

In the context of LDPC code, the decoding energy depends on the required number of iterations and the structure of the parity check matrix [27–31]. Based on message passing, lower bounds on the decoding energy have been derived in [32] and, under an alternate model, in [33].

We introduce a simple model that decomposes the processing tasks in two stages: (1) sampling and (2) decoding. We model the sampler such that each symbol sample has a fixed energy requirement. Thus the energy consumption of the sampler is proportional to the sampling rate, i.e., the fraction of symbol periods in which signal samples are collected. Since the complexity of decoding decreases with the sampling rate, we observe an energy trade-off: we can increase the sampling rate to reduce the decoding energy or collect fewer samples at the expense of additional decoding energy.

The common wisdom, at least among researchers studying decoders, has been that the energy consumption of the sampler is relatively inconsequential compared to that of the decoder. While this may be true, the primary conclusion of this work is that the energy consumption of the sampler, even if it is small, cannot be ignored in an energy harvesting receiver. Depending on the receiver’s finite battery capacity, the per-symbol energy cost of the sampler, relative to the average energy harvesting rate, impacts the reliable communication rate. The key issue is that sampling is a real-time process. When the receiver chooses to sample a packet, the signal samples must be collected during the transmission time of that packet. Thus, the combination of stored battery energy and energy harvested during a slot must be sufficient to ensure the correct sampling rate. By contrast, once the samples have been collected, the decoding can occur offline at a processing rate (and thus energy consumption) matched to the energy harvesting process. This observation leads us to characterize how the battery capacity of the harvesting receiver must grow with the code block length to guarantee reliable communication rates.

## 1.2 Energy Trade-off at the Receiver

Here, we examine energy trade-off ignoring the constraints imposed by the harvesting process or receiver battery size. The receiver must sample the packet and reliably decode it. Due to the randomness of the energy arrivals, the receiving process may be interrupted. This can result in only a subset of transmitted symbols being sampled. On the other hand, a receiver may choose to sample only a subset of symbols in order to save energy for future operations. When the code has block length  $n$  and the sampler recovers samples of  $s$  out of  $n$  symbols, we say the sampling rate is  $s/n$ . We model this selective sampling as an erasure channel concatenated to the original physical channel; symbols that are not sampled are erased. According to [34], if the original physical channel is memoryless with capacity  $C$  and the erasures are independent of the inputs and outputs of that channel and the sampling rate converges in probability to  $\lambda$  or in other words, the proportion of erasures converges in probability to  $1 - \lambda$  (the erasures may have memory), then the capacity of such a channel is  $C_H = \lambda C$ . Fixing a sufficiently long blocklength  $n$  will ensure that a codeword that is sampled at rate  $\lambda$  will be decoded correctly with high probability if  $R < \lambda C$ .

When the transmitter sends packets at rate  $R$  over a channel with capacity  $C_H$ , in various settings, the decoding energy is described through the *capacity gap*,

$$\delta = 1 - \frac{R}{C_H}. \quad (1.1)$$

As an example, in Forney's concatenated codes, decoding energy per channel use,  $\mathcal{E}_D$ , grows exponentially with  $1/\delta$  [35], while in LDPC and Turbo codes, decoding energy scales like  $(1/\delta) \log(1/\delta)$  [27, 28, 30, 31, 36]. Also in polar codes, decoding energy grows polynomially in  $1/\delta$  [37]. Decoding energy, can be viewed as an increasing function of the code rate  $R$  that diverges as  $R$  approaches capacity [32]. It can also be viewed as a decreasing function of the capacity for a fixed code rate  $R$ .

We assume that just  $\lambda$  fraction of the received symbols are sampled, the capacity

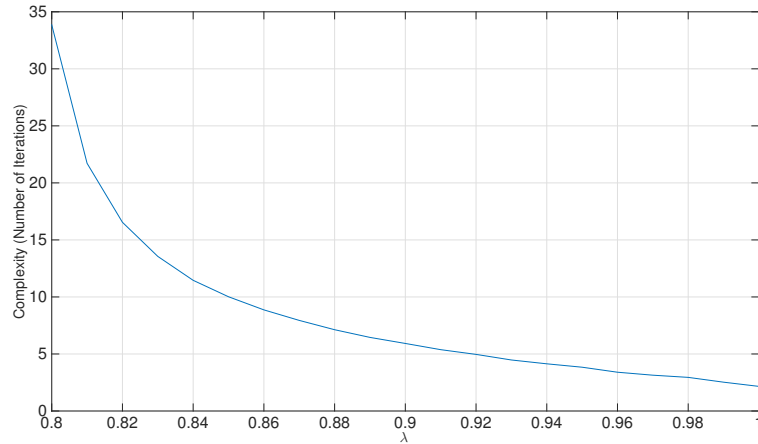


Figure 1.1: Number of iterations versus the sampling rate in the belief propagation iterative algorithm in LDPC code from DVB-S.2 with block length  $n = 64,800$  and code rate  $R = 3/4$  bits/s/Hz.

gap would be

$$\delta = 1 - \frac{R}{\lambda C}, \quad (1.2)$$

and  $\mathcal{E}_D$  is characterized as a function of  $\delta$  the same way as the literature [27–29, 35–38] suggests.

According to the decoding complexity models in [27–29, 35–38], for a fixed  $R$  and  $C$ , the decoding energy is a convex non-increasing function of the normalized code rate  $\lambda$ . To see this, we consider an AWGN channel, with a concatenated erasure channel with erasure rate  $1 - \lambda$ . Fig. 1.1 plots the number of iterations of an LDPC code from DVB-S.2 (Digital Video Broadcasting second generation) standard with  $n = 64,800$  and  $R = 3/4$  bits/sec/Hz. The code is irregular with a specific structure [39]. Fig. 1.2 is based on an Irregular LDPC decoder [40],  $n = 1000$ ,  $R = 1/2$  bits/sec/Hz and the degree distributions are optimized based on the capacity-approaching codes in Richardson et al. [29]. We can see that in both figures, the decoding complexity is a convex non-increasing function. Even if it is not convex, the lower convex-envelope of that can be obtained through time sharing.

We assume  $\nu$  units of energy is expended for sampling. The decoding energy per



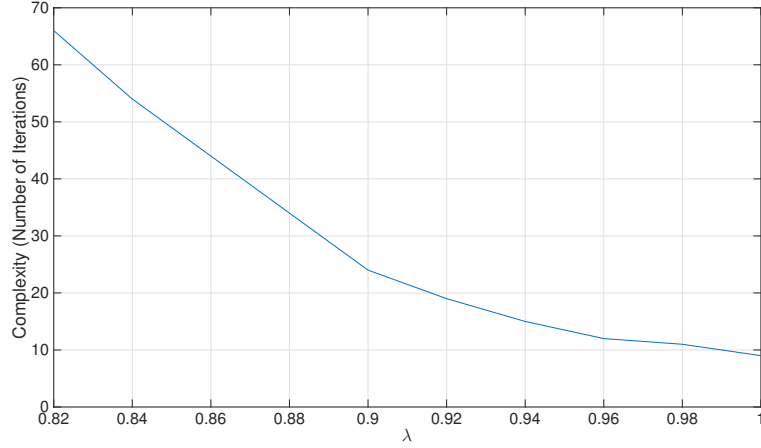


Figure 1.2: Number of iterations versus the sampling rate in the belief propagation iterative algorithm in an Irregular LDPC decoder based on Richardson and Urbanke's 2001 paper,  $n = 1000$ ,  $R = 1/2$  bits/sec/Hz.

symbol consumed by the receiver to reliably decode one packet as a function of the sampling rate  $\lambda$  is denoted by  $\mathcal{E}_D(\lambda)$ .

$$\mathcal{E} = \mathcal{E}(\lambda) = [\nu\lambda + \mathcal{E}_D(\lambda)]. \quad (1.3)$$

Due, to the convexity of  $\mathcal{E}_D(\cdot)$ , for a given  $R$  and  $C$  that there is an optimal sampling rate  $\lambda^* = s^*/n$  such that

$$\lambda^* = \arg \min_{R/C < \lambda \leq 1} \mathcal{E}(\lambda) \quad (1.4)$$

and

$$\mathcal{E}^* = \nu\lambda^* + \mathcal{E}_D(\lambda^*) \quad (1.5)$$

is the minimum energy *per symbol period* required to decode a single rate  $R$  codeword although  $R$  is not included in the notation as we assume it is fixed. Without loss of generality, we can assume  $\nu = 1$ . Furthermore, we emphasize that although energy is expended only on symbols that are sampled,  $\mathcal{E}^*$  amortizes the energy cost of sampling over all symbols, sampled or not. We also note that constraints on the harvesting or on the battery size may preclude the receiver from sampling at the optimal rate  $\lambda^*$ . If

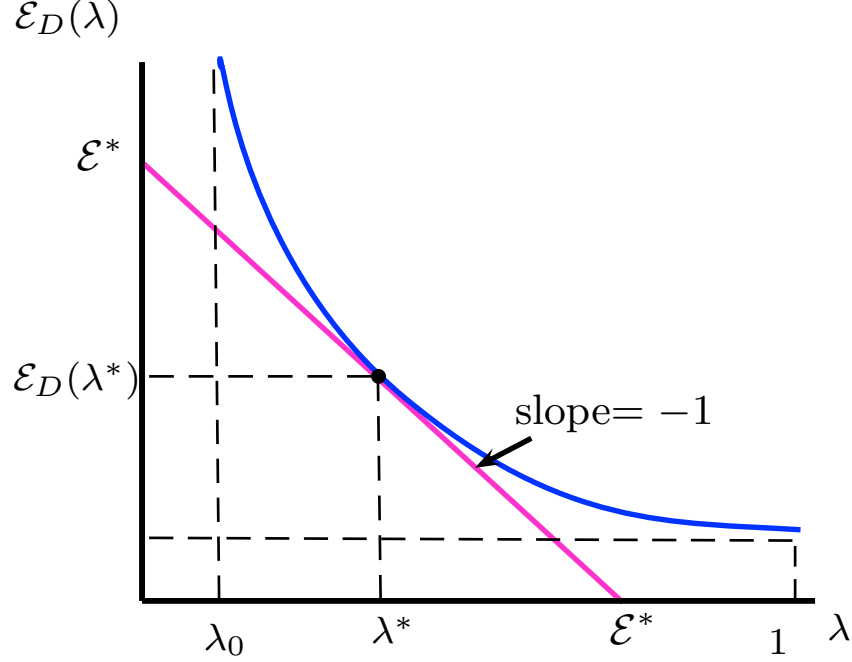


Figure 1.3: Normalized decoding energy  $f(\lambda) = \mathcal{E}_D(\lambda)$  for a fixed  $R$  and  $C$  as function of the sampling rate  $\lambda$ . The total energy per symbol of the receiver is minimized at  $\lambda = \lambda^*$ .

the boundary conditions are not binding, then the minimum in (1.4) occurs when

$$0 = \frac{d\mathcal{E}(\lambda)}{d\lambda} = 1 + \frac{d\mathcal{E}_D(\lambda)}{d\lambda}, \quad (1.6)$$

implying  $d\mathcal{E}_D(\lambda)/d\lambda|_{\lambda=\lambda^*} = -1$ . A geometric representation of this relationship is shown in Figure 1.3. Here, the total energy consumption at each point  $(\lambda, \mathcal{E}_D(\lambda))$  is equal to  $\lambda + \mathcal{E}_D(\lambda)$ . So, if a 45-degree line is plotted crossing that point, it will cut the  $\lambda$ -axis and  $\mathcal{E}_D$ -axis at  $\lambda + \mathcal{E}_D(\lambda)$ . The convexity of  $\mathcal{E}_D(\lambda)$  dictates that the minimum energy is achieved by slightly shifting this 45-degree line with slope  $-1$  until it is tangent to the curve. Note that the sampling rate  $\lambda$  cannot be below  $\lambda_0 = R/C$  as the code rate  $R$  needs to be under the total capacity  $C\lambda$ . At  $\lambda = \lambda_0$ , the gap to the capacity goes to zero and the decoding energy goes to infinity.

We note that this model is consistent with the practical decoding models in [27–29, 31, 35–38]. In addition, we observe that this model does impose restrictions that preclude certain performance enhancements. For example, in a slowly varying channel,

the receiver could exploit channel state information (CSI) to collect its symbol samples when the channel is unusually good. Similarly, the transmitter and receiver could coordinate transmission and reception so that a power-constrained transmitter could use more power for those symbols that the receiver will sample. A coordinated sleep protocol is the limiting case of this approach.

### 1.3 Receiver Sampling and Decoding: Power Comparisons

For receivers, one may ask whether the power consumption of the sampler or the decoder is dominant. Although the power consumption of each system component will necessarily depend on circuit technology, ballpark comparisons that assume the same circuit technology for each can be instructive. Consider first a system in which a 1 Gb/s data stream is encoded at rate 1/2 and then modulated using a 64-QAM constellation. In this case,  $2 \times 10^9/6 \approx 333 \times 10^6$  complex symbols per second are transmitted over an AWGN channel. With bandwidth  $B \approx 333$  MHz, effective number of bits  $\text{ENOB} = 10$ ,  $V_{\text{dd}} = 1.5\text{V}$ , channel length  $L_{\text{min}} = 0.16\mu\text{m}$ , and corner frequency of the  $1/f$  noise  $f_{\text{cor}} = 1$  MHz, the ADC power formula [18, 19]

$$P_{\text{ADC}} \approx \frac{3V_{\text{dd}}^2 L_{\text{min}} (2B + f_{\text{cor}})}{10^{-0.1525\text{ENOB}+4.838}} \quad (1.7)$$

estimates  $P_{\text{ADC}} \approx 228$  mW. According to [41],  $P_{\text{Decoder}} = 690$  mW for the same circuit and system characteristics. Thus,  $P_{\text{Decoder}} \approx 3P_{\text{ADC}}$  in this example system.

We note that the receiver could use a smaller signal constellation, but this will increase the transmission bandwidth, which eventually increases  $P_{\text{ADC}}$ . For example, 4-QAM modulation will raise  $P_{\text{ADC}}$  to 683 mW, which is almost the same as the decoding power.

At lower bit rates, Table 5 in [42] compares different LDPC decoders for IEEE 802.11n applications. For example, code length 1944, 130 nm channel length, code rate 1/2, and throughput of 250 Mb/s yields  $P_{\text{Decoder}} = 76$  mW. Assuming a 16-QAM modulation, the symbol rate is  $2 \times 250/4 = 125\text{M}$  complex symbols per second, implying  $B \approx 125$  MHz. Again using the ADC power formula (1.7),  $P_{\text{ADC}} = 71$  mW.

Thus  $P_{\text{Decoder}} \approx P_{\text{ADC}}$  in this example.

Even though  $P_{\text{ADC}}$  in (1.7) does not include RF front-end power, these simple comparisons suggest that the sampler and decoder have comparable energy consumption and that a receiver energy model should include both system components. Moreover, these simple examples highlight the need for an abstract system model that separates the analysis of rechargeable receivers from technology-specific implementation details.

## 1.4 Notations

We say  $g(n) \in o(n)$  if

$$\lim_{n \rightarrow \infty} \frac{g(n)}{n} = 0. \quad (1.8)$$

The sequence  $X_n$  converges in probability to  $X$ , written  $X_n \xrightarrow{P} X$ , if for every  $\epsilon > 0$ ,

$$\lim_{n \rightarrow \infty} \mathbb{P}[|X_n - X| < \epsilon] = 1. \quad (1.9)$$

$\mathbf{1}(X)$  is an indicator function which is one if the expression  $X$  is true or the event  $X$  occurs and it is zero otherwise.  $[X]^+$  is  $X$  if  $X \geq 0$  and it is zero otherwise.  $\lceil \cdot \rceil$  and  $\lfloor \cdot \rfloor$  also denote the ceiling and the floor of a value, respectively.

## 1.5 Thesis Outline

This work examines how the limited and unreliable source of energy in energy harvesting receivers constrains reliable communication. To model the harvesting receiver, we decompose the processing tasks in two parts: first is sampling or Analog-to-Digital-Conversion (ADC) stage which includes all RF front-end processing, and second is decoding. While sampling and decoding energies are typically comparable, the key issue is that the sampling is a real-time process; the samples must be collected during the transmission time of that packet. Thus the energy harvesting rate and battery size may constrain the sampling rate. In Chapter 2, we propose a system in which, for a

given code rate, channel capacity, and battery size, the receiver can choose the sampling rate to balance the sampling and decoding energy costs and the sender employs a variable-timing codeword transmission scheme that ensures the receiver has the energy resources needed to sample and decode each transmission. We prove the optimality of the combined scheme in a time-invariant channel.

We extend this approach to time-varying fast-fading channels in Chapter 3, where the knowledge of the channel state information is not available at the transmitter. In this case, we will show that channel state knowledge at the receiver can improve the performance of the system. We propose a channel-selective sampling strategy that optimizes a tradeoff between the energy costs of sampling and decoding at the receiver. Based on this tradeoff, we derive a policy maximizing the communication rate and we characterize an energy-constrained rate region.

In a time-varying slow-fading channel, opportunistic transmission when the channel is strong is known to yield a significant gain. A similar result appears in the context of multiuser channels where multiuser diversity gain is achieved by serving users with the strongest channels. In contrast, in Chapter 4, we exploit opportunistic channel selection to reduce the processing energy at the receiver. This saving is derived by increasing the gap between the instantaneous capacity and the code rate, which in turn reduces the required decoding energy. The reduction in processing energy then enables the receiver to sample and process a larger fraction of the packets. We first analyze the single-user case, and propose a signaling scheme that achieves the optimum rate. We then extend this result to a multiuser system.

In Chapter 5 we consider a point-to-point communication channel using Incremental Redundancy Hybrid Automatic Repeat reQuest (IR-HARQ), with limited processing energy. Unlike conventional cases where the processing energy is not limited, here even if the receiver can decode a message, it may still ask the transmitter to send extra redundant bits in order to increase the capacity gap and reduce the processing energy. On the other hand, the extra redundant bits will increase the code length, and may increase the processing energy. To resolve this issue, the receiver can decide to use the extra redundant bits at the receiver, only if it reduces the overall decoding

energy (not just increasing the capacity gap). The other approach is to use Repetition-HARQ scheme, in which the transmitter retransmits the previous packets in response to request for extra redundant bits. This allows the receiver to combine the packets using Maximum Ratio Combining (MRC), and thus keep the effective code length constant. We show that in contrast to the traditional systems, here the Repetition-HARQ could outperform IR-HARQ.

In Chapter 6, we study energy-aware communication under LTE constraints. We consider the problem of energy efficiency in such systems and formulate a problem subject to LTE constraints along with energy constraints. After showing the hardness of the problem, we examine approximation algorithms. However, the classical greedy algorithms do not yield a useful approximation guarantee. So, we propose a deterministic multiplicative updates based algorithm and establish its approximation guarantee.

We will look at two energy efficiency metrics in Chapter 7 and seek to optimize them subject to LTE constraints while assuming that the transmitter is powered by energy harvesting devices dictating energy causality constraints. As the problem is intractable, we seek to optimize a sub-problem and show that it can be reformulated as a constrained submodular set function optimization problem which can be solved approximately using greedy algorithms.

And finally, in Chapter 8, we will present our future directions. They will include working on some extensions to multiuser channels, IR-HARQ schemes and receiver energy optimization in LTE networks.

## Chapter 2

### Energy Harvesting Receivers: Finite Battery Capacity

Harvesting receivers dictate a re-examination of traditional system models. Specifically, receiver design choices influence system performance characteristics that are manifested in both the channel model and the receiver energy consumption. We formulate a model based on a physical channel, a receiver front-end sampler, and a receiver decoder such that channel capacity and sampler energy consumption are jointly specified, independent of the energy trade-off in the receiver between sampling and decoding.

#### 2.1 System Model

We assume a block coding strategy such that a message  $\{1, \dots, 2^{nR}\}$  is communicated by the transmission of a codeword consisting of  $n$  uses of a channel. We will often call a transmitted codeword a *packet* and refer to the transmission period of a codeword as a *slot*. Slots are indexed by  $i = 1, 2, \dots$  such that the codeword transmitted in slot  $i$  is given by the vector  $\mathbf{x}_i = [x_{n(i-1)+1} \ x_{n(i-1)+2} \ \dots, x_{ni}]$  of transmitted symbols. When analyzing just a codeword, the transmitted code word is denoted as  $\mathbf{x} = [x_1 \ x_2 \ \dots, x_n]$  with symbol periods indexed from 1 to  $n$ . The transmitter sends out the next codeword  $\mathbf{x}_{i+1}$  following an idle period of duration  $\tau_i$ . This delay can be chosen such that the receiver is ready to sample and decode the next packet. We call this a *variable-timing* transmission strategy, in contrast to the traditional *fixed-timing* transmission in which  $\tau_i = 0$  so that the end of one slot coincides with the start of the next.

In any event, the receiver front-end processes a symbol or waveform input to produce the symbols  $\mathbf{y} = [y_1 \ \dots \ y_n]$ . We refer to this operation as sampling even though

it may also incorporate demodulation, filtering and quantization. In addition, we refer to the random mapping from  $\mathbf{x}$  to  $\mathbf{y}$  as a physical channel even though elements of the sampling process in the receiver front-end contribute to this mapping. For example, in an AWGN channel, the additive noise in the channel is amplifier noise in the receiver front-end.

We model the energy consumption of the ADC and other RF processing elements in the front-end as requiring  $\nu$  energy units per sample. The constant  $\nu$  is both technology and application dependent. That is, in designing a receiver front-end, the sampling and quantization of the ADC is designed to support the channel bandwidth and SNR needed for communication at intended rates. The receiver front-end design choices are then embodied in the channel from  $\mathbf{x}$  to  $\mathbf{y}$ . We refer to this as the original physical channel and we assume it is memoryless and has single letter capacity  $C$ . As  $C$  depends on the performance of the receiver front-end, which is coupled to the sampling energy  $\nu$ , it is useful to think of  $\nu$  as fixed for a physical channel of capacity  $C$ . Without loss of generality we assume  $\nu = 1$ . That is, energy is measured in the unit of the required energy to take one sample. We note that as the communication rate  $R$  increases, greater precision in sampling is always preferable and often essential, and thus  $\nu$  is a nondecreasing function of the code rate  $R$ . However, for a fixed code rate, the above model is reasonable.

### 2.1.1 Energy Harvesting Model

The energy provided by the environment can be described by a discrete time exogenous stochastic process  $W_t$  of energy arrivals in each symbol period. We assume that the energy arrival,  $W_t$ , is an i.i.d. discrete non-negative random process with PMF  $P_{W_t}$  and mean  $\bar{W}$ . We assume that the energy arrives in integer multiples of the unit sampling energy. In addition, we focus on the case  $\bar{W} < 1$  for which the energy arrival in one symbol period in average is not enough to take a sample. Moreover, we make the technical assumption that the energy arrivals are almost surely bounded. In other words,  $\mathbb{P}[W_t \in [0, b]] = 1$  for some  $b > 0$ . Since a harvesting device cannot harvest energy unboundedly, this constraint is appropriate for the energy harvesting model.



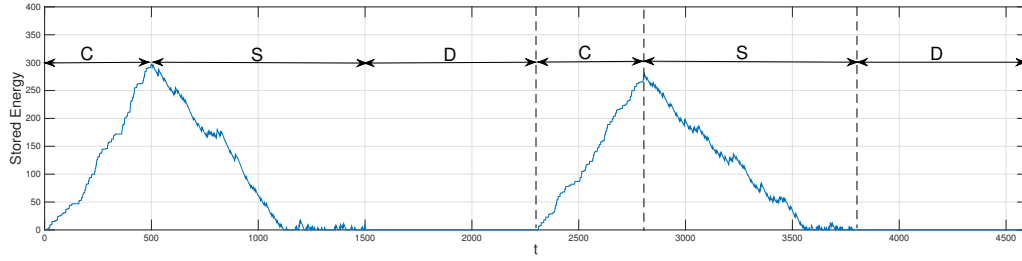


Figure 2.1: Stored energy (in  $\mu J$ ) in different time intervals of variable-timing transmission with random energy harvesting. The intervals marked “C,” “S,” and “D” mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. Here  $W_t = 5 \mu J$  or  $W_t = 1 \mu J$  with probability 0.1 each and  $W_t = 0 \mu J$  with probability 0.8 while  $\bar{W} = 0.6 \mu J$ . Also,  $\beta = 0.3 \mu J$ ,  $n = 1000$  and the sampling rate  $\lambda = 0.9$ .

### 2.1.2 Variable-Timing Transmission

We focus on variable-timing transmission in which after each packet transmission, the transmitter waits for a predetermined time,  $\tau(n)$ , to transmit the next packet. The dependence on the block length  $n$  is necessary since sampling a longer packet may require more energy to be stored in the battery. Also the decoding energy grows with  $n$  as well. During this waiting time interval, the receiver decodes the previously transmitted packet in  $\tau_d(n)$  time units and then, collects energy for sampling the next packet in  $\tau_c(n)$  time units. So, we have

$$\tau(n) = \tau_d(n) + \tau_c(n). \quad (2.1)$$

Note that for the ease of argument we use symbol periods as the time units.

Fig. 2.1 depicts the time intervals for charging, sampling and decoding with random energy arrivals. The optimal decoding and charging time interval is determined by the choice of a sampling rate striking a best tradeoff between sampling and decoding energy in order to maximize the communication rate.

If the harvesting rate or the battery size is large enough, we have the freedom to design a policy that samples packets at the optimal rate  $\lambda^*$  as in (1.4). However,

sampling a packet is an online process that must be completed during the transmission of that packet. Thus the sampling rate may be constrained by the receiver energy that is available in that transmission slot, including both the arrival energy and the stored energy in the battery.

We note that the expected arrival energy in an  $n$ -symbol packet is  $n\bar{W}$ , which scales with  $n$ . In order for the battery to support higher sampling rates, we will see that the battery size must also scale with  $n$ . Thus, we describe an energy harvesting system by a *battery growth rate*  $\beta$  such that we employ a battery of size  $B = \beta n$  when the block length is  $n$ .

With a finite battery, the sampling rate  $\lambda^*$  may not be achievable. In particular, if  $\lambda^* > \bar{W}$ , then harvesting alone will be insufficient to meet the energy requirements for sampling. In this case, pre-charging the battery to  $\beta n$ , along with harvesting energy  $n\bar{W}$  during the packet transmission will enable sampling up to rate  $\beta + \bar{W}$ . However, if  $\beta + \bar{W} < \lambda^*$ , then sampling rate  $\lambda^*$  will not be achievable. We define

$$\tilde{\lambda} = \min \{ \lambda^*, \beta + \bar{W} \}. \quad (2.2)$$

It will be shown later that the sampling rate that maximizes the communication rate will converge to  $\tilde{\lambda}$  for large  $n$ .

### 2.1.3 Performance Metrics

Assume that the transmitter sends packets/codewords  $1, \dots, N(t)$  by time  $t$  with code rate  $R$ . We define the reliable communication rate as the average number of message bits reliably communicated to the receiver per symbol period where “reliably” means that all message bits transmitted are delivered with probability of error that goes to zero as the block length goes to infinity. Accordingly, we define

$$\rho_n(t) = \frac{nR}{t} \sum_{i=1}^{N(t)} I_i^{(n)}, \quad (2.3)$$

where the binary indicator  $I_i^{(n)} = 1$  if the receiver reliably decodes the  $n$ -symbol packet  $i$ . For the overall average communication rate, we define

$$\rho = \lim_{n \rightarrow \infty} \lim_{t \rightarrow \infty} \rho_n(t). \quad (2.4)$$

We seek to find the maximum reliable communication rate, over all feasible policies.

#### 2.1.4 Preliminaries: Chernoff Hoeffding Inequality

Taking  $S$  samples in  $n$  channel uses, a sampling rate of  $\lambda$  is said to be achievable if

$$\lim_{n \rightarrow \infty} \mathbb{P} [S/n < \lambda - n^{\alpha-1}] = 0, \quad \text{for some } \alpha < 1. \quad (2.5)$$

According to Hoeffding's inequality [43], for a set of  $m$  independent random variables  $\{X_1, \dots, X_m\}$ , such that  $\mathbb{P} [X_i \in [0, b]] = 1$  we have

$$\mathbb{P} \left[ \left| \sum_{i=1}^m X_i - \mathbb{E} \left[ \sum_{i=1}^m X_i \right] \right| > d \right] \leq 2 \exp \left( \frac{-2d^2}{mb^2} \right). \quad (2.6)$$

We usually employ the Chernoff-Hoeffding in the following format. For  $m = n$  i.i.d. energy arrivals such that  $\mathbb{P} [W_t \in [0, b]] = 1$ ,

$$\mathbb{P} \left[ \sum_{t=1}^n W_t \leq n\bar{W} - d \right] \leq \exp \left( -\frac{2d^2}{nb^2} \right). \quad (2.7)$$

Note that we will often choose  $d$  as a function of  $n$  such that  $\lim_{n \rightarrow \infty} d/n \rightarrow 0$  while  $\lim_{n \rightarrow \infty} d^2/n \rightarrow \infty$  as in the following lemma for which  $d = n^\alpha$  for some  $\alpha \in (1/2, 1)$ .

**Lemma 2.1** *Let  $\tau(n) = \lceil n\hat{\mathcal{E}}/\bar{W} \rceil$ . For  $1/2 < \alpha < 1$ ,*

$$\mathbb{P} \left[ \sum_{t=1}^{\tau(n)} W_t \leq n\hat{\mathcal{E}} - n^\alpha \right] \leq \exp \left( -\frac{2n^{2\alpha-1}}{b^2(\hat{\mathcal{E}}/\bar{W} + 1/n)} \right).$$

Throughout the rest of this work,  $b > 0$  and  $\alpha \in (1/2, 1)$  are fixed parameters.

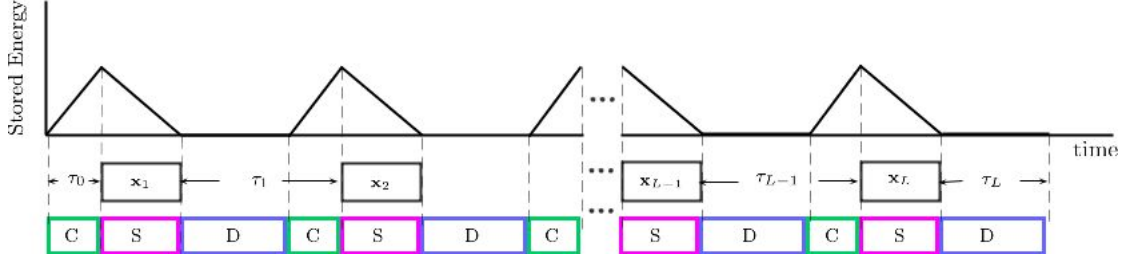


Figure 2.2: Variable-timing optimum policy: Transmitted packets are labeled  $\mathbf{x}_1, \mathbf{x}_2, \dots$  while intervals marked “C”, “S”, and “D” mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. The corresponding graph depicts the receiver’s stored energy.

## 2.2 Achievability: Deterministic Energy Arrivals

When the arrival energy is limited, and especially when the battery size is bounded, the communication rate depends directly on the receiver energy policy. We now describe the charging, sampling and decoding cycle associated with a variable-timing transmission with sampling rate  $\lambda$ . Following [20], we outline the approach for a deterministic energy arrival process in which the energy arrival in any symbol period  $t$  is  $W_t = \bar{W}$  for all  $t$ . We will first analyze the problem under the assumption of  $\lambda > \bar{W}$ , i.e., the deterministic energy arrival rate does not support the intended sampling rate, and then we will go through the case  $\lambda \leq \bar{W}$  where the sampling can be supported by the energy arrival rate. Note that,  $\lambda \leq \lambda_{max} = \min\{\beta + \bar{W}, 1\}$  always holds. Consideration of stochastic energy arrivals is deferred to Section 2.3.

### 2.2.1 Low Harvesting Rate

When  $\bar{W} < \lambda$ , the arrival energy is not large enough to support the sampling. So the receiver needs to pre-charge the battery before the transmission starts. This is due to the fact that sampling is real-time process and as the samples arrive, the sampler need to take them and this process cannot be deferred to a later time when enough energy is available. Variable-timing provides the flexibility to send the next packet when the receiver is ready. As depicted in Fig. 2.2, the transmitter sends packet  $\mathbf{x}_i$  following an idle period of duration  $\tau_{i-1}$  that enables the receiver to decode packet  $\mathbf{x}_{i-1}$  and pre-charge the battery for sampling packet  $\mathbf{x}_i$ . Specifically, as shown in the figure, the

receiver starts by collecting energy  $n(\lambda - \bar{W})$  in time  $\tau_0$ . Next, the transmitter sends packet  $\mathbf{x}_1$  in slot 1 and the receiver samples this packet while also harvesting energy at rate  $\bar{W}$ . However, sampling the packet drains the receiver battery such that the battery is empty at the end of the slot. What follows is a decoding period in which the receiver decodes the sampled packet. The receiver stores no energy in this interval because the decoder is run on a “pay as you go” basis; the decoder runs at a speed such that its energy consumption is matched to the energy harvesting rate  $\bar{W}$ . When decoding of packet  $\mathbf{x}_1$  is completed, the receiver stores energy at rate  $\bar{W}$  in preparation for sampling the next packet. This process of sampling and decoding each packet and pre-charging the battery for sampling the next packet is repeated for each packet. Assuming that all packets are sampled at rate  $\lambda$ , the time required for decoding and pre-charging is the same for all packets (except for the first and last packet).

Denoting the time interval needed for decoding and pre-charging by  $\tau_d(n)$  and  $\tau_c(n)$ , respectively, since energy  $n\mathcal{E}_D(\lambda)$  is collected for decoding the previous packet, we have

$$\tau_d(n) = \left\lceil \frac{n\mathcal{E}_D(\lambda)}{\bar{W}} \right\rceil. \quad (2.8)$$

Then, energy  $n(\lambda - \bar{W})$  is harvested to pre-charge the battery prior to sampling the next packet. When this stored energy is added to the energy  $n\bar{W}$  that is harvested while sampling, the receiver will have sufficient energy to sample at rate  $\lambda$ . Thus,

$$\tau_c(n) = \left\lceil \frac{n(\lambda - \bar{W})}{\bar{W}} \right\rceil. \quad (2.9)$$

Then, the total time between transmissions  $i - 1$  and  $i$  is

$$\begin{aligned} \tau_i(n) &= \tau(n) = \tau_d(n) + \tau_c(n) \\ &= \left\lceil \frac{n\mathcal{E}_D(\lambda)}{\bar{W}} \right\rceil + \left\lceil \frac{n(\lambda - \bar{W})}{\bar{W}} \right\rceil \\ &< \frac{n\mathcal{E}_D(\lambda) + n(\lambda - \bar{W})}{\bar{W}} + 2. \end{aligned} \quad (2.10)$$

And,  $\tau_0(n) = \tau_c(n)$  and  $\tau_L(n) = \tau_d(n)$ .

Assume in a  $[1, T]$  time interval,  $N(T) = L$  packets are transmitted. Assume that using the above algorithm, the receiver decodes all  $L$  packets. So, according to (2.3), and having  $I_i^{(n)} = 1$  for all packets given that  $\tau(n)$  is chosen appropriately, we have

$$\rho_n(T) = \frac{nR}{T} \sum_{i=1}^{N(T)} I_i^{(n)} = \frac{nRL}{T}. \quad (2.11)$$

Since this requires time

$$T = Ln + \sum_{i=0}^L \tau_i = Ln + L\tau(n), \quad (2.12)$$

the throughput is

$$\rho_n(T) = \frac{nRL}{Ln + L\tau(n)} = \frac{nR}{n + \tau(n)} = \rho_n. \quad (2.13)$$

Applying (3.19) to (2.13) and taking  $n$  to infinity, we have

$$\rho = \lim_{n \rightarrow \infty} \rho_n \geq \left( \frac{\bar{W}}{\mathcal{E}_D(\lambda) + \lambda} \right) R. \quad (2.14)$$

### 2.2.2 High Harvesting Rate

When  $\bar{W} \geq \lambda$ , the sampling rate can be supported by the energy arrival and no pre-charging is required. That is, the sampling is finished early along the transmission at time  $\tau_s(n) < n$  and the rest of the transmitted symbols are ignored. We have

$$\tau_s(n) = \left\lceil \frac{n\lambda}{\bar{W}} \right\rceil. \quad (2.15)$$

At time  $\tau_s$ , the receiver has zero stored energy since all harvested energy was used for sampling. Decoding starts immediately after  $\tau_s$  and it takes  $\tau_d(n)$  time units as given by (2.8). If  $\tau_s(n) + \tau_d(n) \leq n$ , no extra time is needed for decoding and the rate  $R$  is not reduced by the process of energy harvesting. Otherwise, time  $\tau_s(n) + \tau_d(n) - n$  is

needed for decoding. So, the time gap is

$$\begin{aligned}
\tau(n) &= [\tau_s(n) + \tau_d(n) - n]^+ \\
&= \left[ \left\lceil \frac{n\lambda}{\bar{W}} \right\rceil + \left\lceil \frac{n\mathcal{E}_D(\lambda)}{\bar{W}} \right\rceil - n \right]^+ \\
&\leq \left[ \frac{n\lambda^* + n\mathcal{E}_D(\lambda)}{\bar{W}} + 2 - n \right]^+.
\end{aligned} \tag{2.16}$$

Substituting (2.16) in (2.13), the communication rate in (2.11) is lower bounded as

$$\rho \geq R \min\left\{1, \frac{\bar{W}}{\mathcal{E}_D(\lambda) + \lambda}\right\}. \tag{2.17}$$

Equation (2.14) and (2.17) suggest that to maximize the communication rate, we need to choose  $\lambda$  to minimize the total required energy  $\lambda + \mathcal{E}_D(\lambda)$ . According to Fig. 1.3, we can see that if there is no bound on the battery size, this minimum will happen at the sampling rate  $\lambda^*$ . However, when the battery is limited, the sampling rate  $\lambda^*$  may not be achievable and  $\beta + \bar{W}$  yields the minimum energy. Thus, the sampling rate at which the energy requirement is minimum is

$$\tilde{\lambda} = \min\{\lambda^*, \beta + \bar{W}\}, \tag{2.18}$$

and then the total minimum energy would be

$$\tilde{\mathcal{E}} = \mathcal{E}_D(\tilde{\lambda}) + \tilde{\lambda}. \tag{2.19}$$

Thus, the maximum achievable rate is

$$\rho = R \min\left\{1, \frac{\bar{W}}{\tilde{\mathcal{E}}}\right\}. \tag{2.20}$$

In general, we must consider the following possibilities:

- $\mathcal{E}_D(\lambda^*) + \lambda^* \leq \bar{W}$ : No extra time is required for pre-charging and decoding. The sampling at rate  $\lambda^*$  and decoding is done before the packet transmission time is over; the achievable throughput is  $R$ .

- $\lambda^* \leq \bar{W}$ : Setting the sampling rate at  $\lambda^*$  to minimize the energy, no pre-charging the battery is needed for sampling as enough energy arrives in each symbol period.
- $\bar{W} < \lambda^* < \beta + \bar{W}$ : The energy collected in one block is not enough to sample at rate  $\lambda^*$ . However, the battery has enough capacity to store such energy. This implies that some time should be spent before sampling to collect energy required for sampling at rate  $\lambda^*$ .
- $\beta + \bar{W} < \lambda^*$ : Not only is the energy collected in one slot insufficient for sampling at rate  $\lambda^*$ , the battery is also not big enough to enable sampling at this rate. By fully charging the battery prior to sampling, the largest possible sampling rate is  $\beta + \bar{W}$ . In this case, the convexity of the decoding energy function, as shown in Fig. 1.3, dictates that the minimum total energy occurs at the maximum sampling rate. Therefore, to maximize the communication rate, we should set the sampling rate at  $\beta + \bar{W}$ .

As a result, the following theorem is concluded.

**Theorem 2.1** *A variable-timing transmission system with packets encoded at rate  $R$  while the receiver is harvesting energy deterministically at rate  $\bar{W}$  can achieve the communication rate*

$$\rho = R \min\left\{1, \frac{\bar{W}}{\bar{\mathcal{E}}}\right\}.$$

In the sequel, we show that Theorem 2.1 holds when the energy arrival process  $W_t$  is an i.i.d. random process with  $\mathbb{E}[W_t] = \bar{W}$ . With a stochastic energy arrival process, the time required to pre-charge the battery, the number of symbol samples collected, and the time required to harvest the energy to run the decoder may all be random. However, as the block length  $n$  grows, the law of large numbers will prevail and these variations will be relatively insignificant.

### 2.3 Achievability: Random Energy Arrivals

In the stochastic case, we verify Theorem 2.1 in three steps:



- Starting from an empty battery, we first charge the battery such that the battery reaches the target level  $n\mu$  such that

$$\mu = \min \left\{ [\lambda^* - \bar{W}]^+, \beta \right\}. \quad (2.21)$$

This requires charging the battery for time

$$\tau_c(n) = \left\lceil \frac{n\mu}{\bar{W}} \right\rceil. \quad (2.22)$$

We note that if  $\lambda^* < \bar{W}$ , then no pre-charging is required.

- With the battery charged, sampling of the next transmitted packet occurs at the optimal rate  $\tilde{\lambda}$ .
- Finally, decoding is done while energy harvesting continues. When  $\lambda^* > \bar{W}$ , the sampling is in progress until the end of the slot and decoding starts after the packet transmission time is over. The decoding time is set to

$$\tau_d(n) = \left\lceil \frac{n\mathcal{E}_D(\tilde{\lambda}) + n^\alpha}{\bar{W}} \right\rceil. \quad (2.23)$$

However, when  $\lambda^* \leq \bar{W}$ , the sampling is done early in the middle of the packet transmission and it is not efficient to wait until the end of the slot to start decoding. In this case, the decoding is started immediately after the samples are taken.

In the following subsections, we describe these three steps when  $\lambda^* > \bar{W}$ . We will explain the other case in Section 2.3.2.

### 2.3.1 Low Harvesting Rate

We first analyze the problem assuming  $\bar{W} < \lambda^*$ .

### Pre-charging the battery

Here, charging the battery for time  $\tau_c(n)$  implies that we harvest energy

$$U_0 = \sum_{t=1}^{\tau_c(n)} W_t. \quad (2.24)$$

The concentration inequality in Lemma 2.1 yields

$$\mathbb{P}[U_0 \leq n\mu - n^\alpha] \leq \exp \left[ -\frac{2n^{2\alpha-1}}{b^2(\mu/\bar{W} + 1/n)} \right]. \quad (2.25)$$

That is, for large  $n$ , the initial battery storage is asymptotically close to  $n\mu$  with high probability.

### Packet Sampling

Starting the sampling period with energy  $U_0$ , the stored energy at the start of symbol period  $t$  is denoted by  $U_{t-1}$ . The energy  $W_{t-1}$  harvested in symbol period  $t-1$  is available for sampling symbol  $t$ . If  $U_{t-1} + W_{t-1} \geq 1$ , then the symbol  $t$  is sampled. The stored energy at time  $t$  is

$$U_t = \min \{ n\beta, [U_{t-1} + W_{t-1} - 1]^+ \}. \quad (2.26)$$

During the sampling interval, harvested energy is used only for sampling; each sample requires one unit of energy. Thus, by energy conservation, the number of samples collected is

$$S = U_0 + \sum_{t=1}^{n-1} W_t - V - U_n, \quad (2.27)$$

where  $V$  denotes the energy that is discarded over the sampling period because the battery is full and  $U_n$  is the extra energy left in the battery at the end of the time slot.

We now prove that the sampling rate  $\tilde{\lambda}$  in (2.18) is achievable. Fig. 2.3 shows a sample path of the stored energy,  $U_t$  over a sampling interval. According to this figure, the arrival energy  $W_t > 0$ , can be viewed as a job with service time  $W_t$  arriving at an

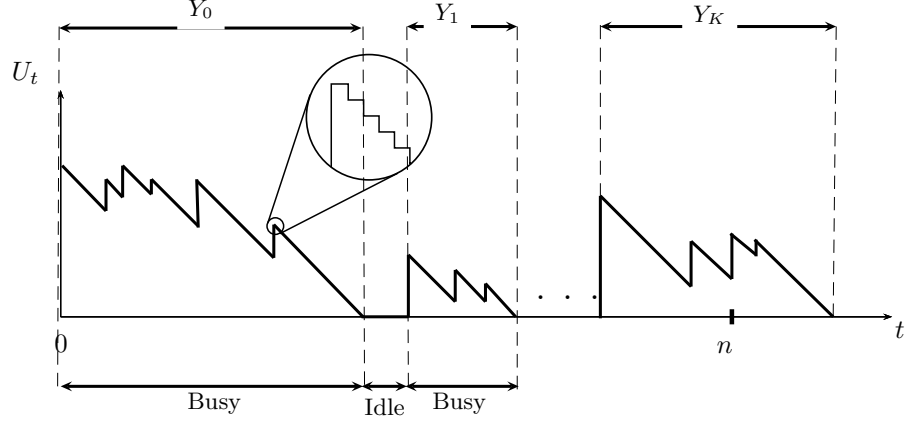


Figure 2.3: Stored energy over the sampling interval modeled as the energy queue. Comparing to Fig. 2.1 and Fig. 2.2, here just a single (S)ampling interval is shown.

energy queue at time  $t$ . The queue service time distribution associated with this job is given by  $P_{W_t|W_t>0}$ . As the energy arrivals are i.i.d., the inter-arrival time between two consecutive job arrivals is a geometric random variable with expected value  $1/\mathbb{P}[W > 0]$ . Thus the “jobs” are arriving as a memoryless random process and the energy queue is considered as discrete-time M/G/1.

We observe that excess energy due to limited battery is discarded only during *busy periods* in which  $U_t \geq 1$ . The busy period  $i$  starts at time  $t_i$  if  $U_{t_i} \geq 1$  while  $U_{t_i-1} = 0$ . This busy period is over the first time  $t > t_i$  such that  $U_t = 0$ .

We assume that there are  $K+1$  busy periods  $0, 1, \dots, K$  over a time slot. Busy period 0 starts with energy  $U_0$ . Busy periods  $i, i = 1, 2, \dots, K$  start from zero initial energy. We use  $V_i$  and  $Y_i$  to denote the energy discarded in busy period  $i$  and the length of busy period  $i$ , respectively. Assuming that the sampling continues indefinitely,  $V_1, V_2, \dots$  are i.i.d.. Note however that if  $U_n \geq 1$ , the packet ends while the busy period  $K$  is still in progress. In this case, we assume a model such that energy arrival, energy harvesting and sampling continue to the end of the busy period past time  $n$  while the samples taken after time  $n$  are ignored. Under this model, the energy  $V$  discarded during the block satisfies

$$V = \sum_{i=0}^{K-1} V_i + V'_K \leq \sum_{i=0}^K V_i, \quad (2.28)$$

where  $V'_K$  is the energy loss in busy period  $K$  by time  $n$  due to the limited battery.

### Achievable Sampling Rate

According to (2.27), there are two components of energy loss. First is the residual energy in the battery at time  $n$ ,  $U_n$  which is upper bounded by  $Y_K$ . We will show that this energy is negligible. The second component of energy waste is the new arrival energy that is discarded because the battery is full. We will show that this is also negligible.

To show the former, we note that knowledge that  $K = k$  can influence the conditional distribution of  $V_1, \dots, V_k$ . Specifically, when  $K$  is large the busy periods are short and the likelihood of energy loss is small. In addition, the length of a busy period with unbounded battery is an upper bound to  $Y_K$ . Using these facts, we prove the following lemma.

#### Lemma 2.2

$$\mathbb{E}[Y_i] \leq \frac{\mathbb{E}[W|W > 0]}{1 - \bar{W}}, \quad i = 1, 2, \dots, K - 1$$

and

$$\mathbb{E}[Y_K] \leq \eta_W \equiv \frac{\mathbb{E}[W^2|W > 0]}{(1 - \bar{W})^2 \mathbb{E}[W|W > 0]}.$$

#### Proof:

Assume that a sample point  $\Theta$  in the sample space of the energy arrivals maps into a sequence of energy arrivals  $\{w_t\}$ . Then  $\{u_t^i\}$  is the sample path of the stored energy (starting from zero energy) in busy period  $i$  and  $y_i$  is the length of this busy period while sampling is done when the battery energy is at least one unit. Assume that  $\{\tilde{u}_t^i\}$  is the sequence of the stored energy corresponding to the same sample point  $\Theta$  starting from zero energy when the battery is unconstrained.  $\tilde{y}_i$  is also the length of the busy period corresponding to  $\{\tilde{u}_t^i\}$ . Assume at  $t_1$ ,  $u_t^i$  and  $\tilde{u}_t^i$  hit  $B$ . Note that  $u_t^i = \tilde{u}_t^i$  for  $t \leq t_1$  as both sequences start from zero energy and experience the same energy arrival while sampling in every symbol period. But after  $t_1$ ,  $\tilde{u}_t$  may get larger than  $u_t$  if the

arrival energy is larger than 1. From that point on,  $\tilde{u}_t^i$  is never less than  $u_t^i$ . So,

$$\begin{cases} \tilde{u}_t^i = u_t^i, & t \leq t_1 \\ \tilde{u}_t^i \geq u_t^i, & t > t_1. \end{cases} \quad (2.29)$$

So, (2.29) implies that  $\tilde{u}_t^i$  crosses zero no earlier than  $u_t^i$ , so

$$y_i \leq \tilde{y}_i. \quad (2.30)$$

As this is true for all sample paths  $\Theta$  in sample space, so we have

$$\mathbb{E}[Y_i] \leq \mathbb{E}[\tilde{Y}_i]. \quad (2.31)$$

According to [44], the mean of the busy period  $\tilde{Y}$  for an M/G/1 queue starting from a zero queue when there is no limit on the size of the queue (unbounded battery) is

$$\mathbb{E}[\tilde{Y}_i] = \frac{\mathbb{E}[W|W > 0]}{1 - \bar{W}}. \quad (2.32)$$

So, (2.31) and (2.32) imply

$$\mathbb{E}[Y_i] \leq \frac{\mathbb{E}[W|W > 0]}{1 - \bar{W}}, \quad i = 1, 2, \dots, K-1. \quad (2.33)$$

Note that  $\mathbb{E}[Y_i]$  is limited and doesn't grow with  $n$ . We use  $Y$  and  $\tilde{Y}$  as the random variables identical to  $Y_i$  and  $\tilde{Y}_i$ ,  $i = 1, 2, \dots, K$  respectively.

On the other hand, as the block is terminated in the busy period  $K$ , the random incidence effect suggests that the length of this busy period be larger than the other ones, therefore, according to Section 2.13 in [45],

$$f_{\tilde{Y}_K}(\tilde{y}) = \frac{\tilde{y} f_{\tilde{y}}(\tilde{y})}{\mathbb{E}[\tilde{Y}]}, \quad (2.34)$$

where the coefficient  $1/\mathbb{E}[\tilde{Y}]$  is derived from the fact that the pdf must integrate to 1.

Then, taking the expectation of  $\tilde{Y}_K$  yields

$$\mathbb{E}[\tilde{Y}_K] = \frac{\mathbb{E}[\tilde{Y}^2]}{\mathbb{E}[\tilde{Y}]}. \quad (2.35)$$

According to Section 5.8 in [44], where different moments of the busy period of an  $M/G/1$  queue are derived using the Laplace transform, we have

$$\mathbb{E}[\tilde{Y}^2] = \frac{\mathbb{E}[W^2|W > 0]}{(1 - \bar{W})^3}. \quad (2.36)$$

It follows from (2.32), (2.35) and (2.36) that

$$\mathbb{E}[\tilde{Y}_K] = \frac{\mathbb{E}[W^2|W > 0]}{(1 - \bar{W})^2 \mathbb{E}[W|W > 0]}. \quad (2.37)$$

Thus the claim follows (2.31) and (2.37).  $\square$

Using the Markov inequality,

$$\mathbb{P}[Y_K \geq n^\alpha] \leq \frac{\eta W}{n^\alpha} \quad (2.38)$$

Obviously, as  $U_n \leq Y_K$ , we have

$$\mathbb{P}[U_n \geq n^\alpha] \leq \frac{\eta W}{n^\alpha}. \quad (2.39)$$

Thus, the leftover energy in the battery at the end of the block, is not arbitrarily large with a high probability. Now, we show that the energy loss  $V$  due to the limited battery is not significant. Assume the Moment Generating Function (MGF) of  $W$ , which is defined as  $g_W(r) = \mathbb{E}[\exp(rW)]$ , exists (i.e., is finite) in an interval  $(r_-, r_+)$ , such that  $r_- < 0$  and  $r_+ > 0$ . Defining  $X = W - 1$ , we define the semi-invariant MGF [46] of  $X$  as

$$\gamma_X(r) \triangleq \ln g_X(r) = \ln \mathbb{E}[\exp(rX)] \quad (2.40)$$

Recalling that we have assumed  $\bar{W} < 1$ , it follows that  $\mathbb{E}[X] < 0$ .

**Lemma 2.3** *There exists  $r^* > 0$  where  $\gamma_X(r^*) = 0$  such that for some  $h(n)$*

$$\mathbb{P}[V_0 > 0 | U_0 = n\beta - h(n) \geq 0] \leq \exp(-r^*h(n)).$$

**Proof:** Here, we focus on the first busy period. Defining  $X_t = W_t - 1$ , assuming that during a busy period the samples are taken, we have the following iterative relationship for the stored energy at the battery at the end of symbol period  $t$

$$U_t = \max\{U_{t-1} + X_{t-1}, 0\}. \quad (2.41)$$

Assuming the first busy period ends at time  $t_1$ , it follows

$$\mathbb{P}[U_t \geq B] \geq \mathbb{P}[V_0 > 0] \quad \forall t \in [0, t_1] \quad (2.42)$$

So, it is sufficient to derive the probability of  $U_t$  hitting  $B$  while the stored energy starts from the energy level  $U_0$ . Note that we assume  $W_0 = 0$  so,  $X_0 = -1$ . We have

$$\begin{aligned} U_t &= \max\{U_{t-1} + X_{t-1}, 0\} \\ &= \max\{U_{t-2} + X_{t-2} + X_{t-1}, X_{t-1}, 0\} \\ &\dots \\ &= \max\{(X_0 + \dots + X_{t-1}), (X_1 + \dots + X_{t-1}), \dots, (X_{t-2} + X_{t-1}), X_{t-1}, 0\}. \end{aligned} \quad (2.43)$$

So for some  $a$ ,

$$\begin{aligned} \mathbb{P}[U_t \geq a] \\ &= \mathbb{P}[\max\{(X_0 + \dots + X_{t-2} + X_{t-1}), \dots, (X_{t-2} + X_{t-1}), X_{t-1}, 0\} \geq a]. \end{aligned} \quad (2.44)$$

This probability is equal to the probability that a random walk based on  $X_t$  crosses  $a$  by the trial  $t$ . According to [46], it can be shown that (2.44) implies that

$$\mathbb{P}[U_t \geq a] \leq \mathbb{P}[U_{t+1} \geq a]$$

And since this sequence is increasing and it can not exceed 1, it has a limit when  $t \rightarrow \infty$  which is denoted as  $\mathbb{P}[U \geq a]$ .

We are interested to know the probability of  $U_t$  crossing  $B$  for the first time before crossing 0. Given  $U_0 = B - h(n)$ , this is equivalent to the event that  $\hat{U}_t = U_t - B + h(n)$ , for some  $h(n)$ , crosses  $h(n)$  before crossing  $h(n) - B$ . As  $\mathbb{E}[X] < 0$ , according to Corollary 9.4.1 in [46], the Wald's identity for two thresholds will result in the following exponential bound

$$\mathbb{P}[U_J \geq B] = \mathbb{P}[\hat{U}_J \geq h(n)] \leq \exp(-r^* h(n)), \quad (2.45)$$

where  $r^*$  is chosen such that  $\gamma(r^*) = 0$ , while  $\gamma(r) = \gamma_X(r) = \ln \mathbb{E}[\exp(rX)]$  (semi-invariant MGF). To prove that such an  $r^* > 0$  exists, we define  $g_X(r) = \mathbb{E}[\exp(rX)]$ . Then, we have

$$g_X(r) = \mathbb{E}[\exp(rX)] \quad (2.46a)$$

$$\frac{dg_X(r)}{dr} = g'_X(r) = \mathbb{E}[X \exp(rX)] \quad (2.46b)$$

$$\frac{d^2 g_X(r)}{dr^2} = g''_X(r) = \mathbb{E}[X^2 \exp(rX)]. \quad (2.46c)$$

Also, as  $\gamma_X(r) = \ln g_X(r)$ ,

$$\gamma'_X(r) = \frac{\mathbb{E}[X \exp(rX)]}{\mathbb{E}[\exp(rX)]} = \frac{g'_X(r)}{g_X(r)}, \quad (2.47)$$

so,

$$\gamma'_X(0) = \mathbb{E}[X] < 0, \quad (2.48)$$

and

$$\begin{aligned} \gamma''(r) &= \frac{\mathbb{E}[X^2 \exp(rX)] \mathbb{E}[\exp(rX)] - (\mathbb{E}[X \exp(rX)])^2}{(\mathbb{E}[\exp(rX)])^2} \\ &= \frac{g''_X(r) g_X(r) - g'_X(r)^2}{g_X(r)^2}. \end{aligned} \quad (2.49)$$



On the other hand, according to exercise 1.26 in [46], if the MGF of  $W$  exists on  $I(W)$ , the MGF of  $X = W - 1$  also exists on that interval. Also, it can be seen that  $g''_{X-c}(r) \geq 0$ . As suggested in that exercise, choosing  $c = g'_X(r)/g_X(r)$  will result in  $g''_X(r)g_X(r) - g'_X(r)^2 \geq 0$ . It then follows from (2.49) that  $\gamma''_X(r) \geq 0$ . So  $\gamma_X(r)$  is convex in  $(r_-, r_+)$  and considering (2.48) and  $\gamma_X(0) = 0$ , we conclude that  $\gamma_X(r)$  cuts the r-axis in the positive region. Therefore, such an  $r^* > 0$  as in (2.45) exists.  $\square$

If  $\lim_{n \rightarrow \infty} h(n) = \infty$ , then as  $n \rightarrow \infty$  the above probability goes to zero. As a result, in order to avoid energy loss in busy period zero, the pre-charging interval is set such that  $\lim_{n \rightarrow \infty} h(n) = \infty$ .

Lemma 2.3 can easily be extended to other busy periods. Knowing that other busy periods start from zero stored energy, that is  $h(n) = n\beta$ , the following corollary holds.

**Corollary 2.1** *There exists  $r^* > 0$  satisfying  $\gamma_X(r^*) = 0$  such that*

$$\mathbb{P}[V_i > 0] \leq \exp(-n\beta r^*) \quad \forall i \geq 1.$$

**Lemma 2.4** *For random variables  $A_1, A_2, \dots, A_m$*

$$\mathbb{P}\left[\sum_{i=1}^m A_i \leq 0\right] \leq \sum_{i=1}^m \mathbb{P}[A_i \leq 0].$$

**Proof:** It is easy to see this for two random variables  $A_1$  and  $A_2$ . We have

$$\{A_1 + A_2 \leq 0\} \subseteq \{A_1 \leq 0\} \cup \{A_2 \leq 0\}, \quad (2.50)$$

then

$$\begin{aligned} \mathbb{P}[A_1 + A_2 \leq 0] &\leq \mathbb{P}[A_1 \leq 0 \cup A_2 \leq 0], \\ &\leq \mathbb{P}[A_1 \leq 0] + \mathbb{P}[A_2 \leq 0] \end{aligned} \quad (2.51)$$

where here we used the general union bound. The same is proved for  $m$  random variables using induction.  $\square$

Now using Lemma 2.2, Lemma 2.3, Lemma 2.4, and Corollary 2.1, we prove the

following theorem.

**Theorem 2.2** *If  $\lambda^* > \bar{W}$ ,*

$$\mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \leq \epsilon(n),$$

where  $\epsilon(n) = \exp(-r^*h(n)) + n \exp(-n\beta r^*) + 2n^{-\alpha}\eta_W + \exp\left(-\frac{2(n^\alpha/2 - \bar{W})^2}{(n-1)b^2}\right)$  while  $h(n) = n^\alpha$  if  $\lambda^* \geq \bar{W} + \beta$  and  $h(n) = n\beta - n\lambda^* + n\bar{W} + n^\alpha$  otherwise. Note that  $\lim_{n \rightarrow \infty} \epsilon(n) = 0$  in both cases.

**Proof:** Using (2.27) and (2.28), the number of samples  $S$  satisfies

$$S \geq U_0 + \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n. \quad (2.52)$$

This implies

$$\begin{aligned} & \mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \\ & \leq \mathbb{P} \left[ U_0 + \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n\tilde{\lambda} - 2n^\alpha | U_0 = n\mu - n^\alpha \right] \\ & = \mathbb{P} \left[ n \min\{\lambda^* - \bar{W}, \beta\} - n^\alpha + \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n \min\{\beta + \bar{W}, \lambda^*\} - 2n^\alpha \right]. \end{aligned}$$

If  $\lambda^* - \bar{W} \geq \beta$ , then  $U_0 = n\beta - n^\alpha$  and  $\tilde{\lambda} = \beta + \bar{W}$ . The above expression then reduces to

$$\begin{aligned} & \mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \\ & \leq \mathbb{P} \left[ n\beta - n^\alpha + \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n\beta + n\bar{W} - 2n^\alpha \right] \\ & = \mathbb{P} \left[ \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n\bar{W} - n^\alpha \right]. \end{aligned} \quad (2.53)$$

Since  $K \leq n$ , we have

$$\sum_{i=1}^K V_i \leq \sum_{i=1}^n V_i, \quad (2.54)$$

Eq. (2.53) is equivalent to

$$\mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \leq \mathbb{P} \left[ \sum_{t=1}^{n-1} W_t - V_0 - \sum_{i=1}^n V_i - U_n < n\bar{W} - n^\alpha \right]. \quad (2.55)$$

We define the random variables  $A_i$ ,  $i = 0, 1, \dots, n+2$  as

$$A_0 = -V_0 \quad (2.56a)$$

$$A_i = -V_i, \quad i = 1, 2, \dots, n \quad (2.56b)$$

$$A_{n+1} = -U_n + n^\alpha/2 \quad (2.56c)$$

$$A_{n+2} = \sum_{t=1}^{n-1} W_t - (n-1)\bar{W} + n^\alpha/2 - \bar{W}. \quad (2.56d)$$

Following Lemma 2.3, we choose  $h(n) = n^\alpha$ ,  $1/2 < \alpha < 1$ . So, we have

$$\mathbb{P} [A_0 < 0] = \mathbb{P} [V_0 > 0] \leq \exp(-n^\alpha r^*). \quad (2.57)$$

where the upper bound goes to zero as  $n \rightarrow \infty$ . For the consequent busy periods,  $i \geq 1$ , following Corollary 2.1,

$$\mathbb{P} [A_i < 0] = \mathbb{P} [V_i > 0] \leq \exp(-nr^*\beta), \quad i = 1, 2, \dots, K. \quad (2.58)$$

Then, (2.55) is equivalent to

$$\begin{aligned} \mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \\ \leq \mathbb{P} [A_0 + A_1 + \dots A_{n+2} < 0] \end{aligned} \quad (2.59)$$

$$\stackrel{(a)}{\leq} \mathbb{P} [A_0 < 0] + \sum_{i=1}^n \mathbb{P} [A_i < 0] + \mathbb{P} [A_{n+1} \leq 0] + \mathbb{P} [A_{n+2} \leq 0] \quad (2.60)$$

$$\leq \epsilon(n). \quad (2.61)$$

Note that we have used Lemma 2.4 in (a).  $\mathbb{P} [A_{n+1} \leq 0]$  and  $\mathbb{P} [A_{n+2} \leq 0]$  are also upper bounded using (2.39) and (2.7), respectively. Note that the bound,  $\epsilon(n)$ , goes to zero as  $n$  goes to infinity.

If  $\beta > \lambda^* - \bar{W}$ , then  $U_0 = n\lambda^* - n\bar{W} - n^\alpha$  and  $\tilde{\lambda} = \lambda^*$ . Then,

$$\begin{aligned} \mathbb{P} \left[ S/n < \tilde{\lambda} - 2n^{\alpha-1} | U_0 = n\mu - n^\alpha \right] \\ \leq \mathbb{P} \left[ n\lambda^* - n\bar{W} - n^\alpha + \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n\lambda^* - 2n^\alpha \right] \end{aligned} \quad (2.62a)$$

$$= \mathbb{P} \left[ \sum_{t=1}^{n-1} W_t - \sum_{i=0}^K V_i - U_n < n\bar{W} - n^\alpha \right], \quad (2.62b)$$

which is the same expression as (2.53). Note that in this case,  $h(n) = n\beta - n\lambda^* + n\bar{W} + n^\alpha$ . Therefore, according to Lemma 2.3 we have

$$\mathbb{P} [V_0 > 0] \leq \exp(-r^*h(n)), \quad (2.63)$$

where as  $\lim_{n \rightarrow \infty} h(n) = \infty$ , the above probability goes to zero for large  $n$ . Following the same steps as in (2.60) through (2.61), the claim is proved in this case as well.

□

### Achievable Communication Rate

Choosing the decoding time interval,  $\tau_d(n)$  as in (2.23) ensures that  $\mathcal{E}_D(\tilde{\lambda})$  energy arrives in this time interval with high probability. Using Lemma 2.1,

$$\mathbb{P} \left[ \sum_{t=1}^{\tau_d(n)-1} W_t \leq n\mathcal{E}_D(\tilde{\lambda}) \right] \leq \exp \left( -\frac{2n^{2\alpha-1}}{b^2((\mathcal{E}_D(\tilde{\lambda}) + n^{\alpha-1})/\bar{W} + 1/n)} \right). \quad (2.64)$$

We assume that a packet is decoded if the sampling rate  $\tilde{\lambda}$  is achieved but is discarded otherwise. Decoding cannot be done partially and it needs to be done fully after sampling. The decoding is completed only if after  $\tau_d$  symbol periods, at least  $n\mathcal{E}_D(\tilde{\lambda})$  energy units are harvested. Otherwise, the packet is discarded.

In general, failure in decoding packet  $i$  can be due to an insufficient number of samples or insufficient energy for decoding. Insufficient samples can be caused by low energy in the charging period such that the accumulated energy over the  $\tau_c$  interval is not enough to achieve  $\tilde{\lambda}$ . It also can be the result of low energy arrival during the sampling period; here, the initial energy may be enough but still insufficient energy arrival during the transmission may lower the sampling rate.

$$\begin{aligned} \mathbb{P} [S < n\tilde{\lambda} - 2n^\alpha] &= \mathbb{P} [S \leq n\tilde{\lambda} - 2n^\alpha | U_0 \geq n\mu - n^\alpha] \mathbb{P} [U_0 \geq n\mu - n^\alpha] \\ &\quad + \mathbb{P} [S < n\tilde{\lambda} - 2n^\alpha | U_0 < n\mu - n^\alpha] \mathbb{P} [U_0 < n\mu - n^\alpha] \end{aligned} \quad (2.65a)$$

$$\leq \mathbb{P} [S < n\tilde{\lambda} - 2n^\alpha | U_0 = n\mu - n^\alpha] + \mathbb{P} [U_0 < n\mu - n^\alpha] \quad (2.65b)$$

$$\leq \epsilon(n) + \exp \left[ -\frac{2n^{2\alpha-1}}{b^2(\mu/\bar{W} + 1/n)} \right], \quad (2.65c)$$

where  $\epsilon(n)$  is defined in Theorem 2.2.  $\mathbb{P} [U_0 \geq n\mu - n^\alpha]$  and  $\mathbb{P} [S < n\tilde{\lambda} - 2n^\alpha | U_0 < n\mu]$  in (2.65a) are both upper bounded by 1. Note that  $\mathbb{P} [S \leq n\tilde{\lambda} - 2n^\alpha | U_0 \geq n\mu - n^\alpha]$  is also upper bounded by  $\mathbb{P} [S \leq n\tilde{\lambda} - 2n^\alpha | U_0 = n\mu - n^\alpha]$ . That is, as the initial stored energy goes up, the probability of taking samples less than a threshold goes down. Eq. (2.65b) follows from the upper bounds in Theorem 2.2 and (2.25) respectively. The

probability that packet  $i$  fails to be decoded satisfies

$$\mathbb{P} \left[ I_i^{(n)} = 0 \right] = \mathbb{P} \left[ \{S < n\tilde{\lambda} - 2n^\alpha\} \cup \left\{ \sum_{t=1}^{\tau_d-1} W_t < n\mathcal{E}_D(\tilde{\lambda}) \right\} \right] \quad (2.66a)$$

$$\leq \mathbb{P} \left[ S < n\tilde{\lambda} - 2n^\alpha \right] + \mathbb{P} \left[ \sum_{t=1}^{\tau_d-1} W_t \leq n\mathcal{E}_D(\tilde{\lambda}) \right] \quad (2.66b)$$

$$\leq \epsilon(n) + \exp \left[ -\frac{2n^{2\alpha-1}}{b^2(\mu/\bar{W} + 1/n)} \right] + \exp \left( -\frac{2n^{2\alpha-1}}{b^2((\mathcal{E}_D(\tilde{\lambda}) + n^{\alpha-1})/\bar{W} + 1/n)} \right), \quad (2.66c)$$

where (2.66b) follows from the union bound and (2.66c) follows from the upper bound (2.65c) and (2.64). As a result,

$$\lim_{n \rightarrow \infty} \mathbb{P} \left[ I_i^{(n)} = 1 \right] = 1. \quad (2.67)$$

We assume that each episode of pre-charging, sampling and decoding starts with zero energy at the battery to facilitate the analysis. In particular, if there is any energy left from previous episode it is discarded. As a result, a renewal occurs at the start of each episode.

According to (2.3), the communication rate is

$$\begin{aligned} \rho_n &= \lim_{t \rightarrow \infty} \rho_n(t) \\ &= \lim_{t \rightarrow \infty} \frac{nR}{t} \sum_{i=1}^{N(t)} I_i^{(n)} \\ &= \left( \lim_{t \rightarrow \infty} \frac{nRN(t)}{t} \right) \left( \lim_{t \rightarrow \infty} \frac{1}{N(t)} \sum_{i=1}^{N(t)} I_i^{(n)} \right). \end{aligned} \quad (2.68)$$

Since  $N(t) = \left\lfloor \frac{t}{n + \tau(n)} \right\rfloor$ ,

$$\lim_{t \rightarrow \infty} \frac{N(t)}{t} = \frac{1}{n + \tau(n)}. \quad (2.69)$$

In addition, since  $\lim_{t \rightarrow \infty} N(t) = \infty$  and the  $I_i^{(n)}$  are i.i.d., the law of large numbers

implies

$$\lim_{t \rightarrow \infty} \frac{1}{N(t)} \sum_{i=1}^{N(t)} I_i^{(n)} = \mathbb{P} \left[ I_i^{(n)} = 1 \right]. \quad (2.70)$$

Then,

$$\rho_n = \left( \frac{nR}{n + \tau(n)} \right) \mathbb{P} \left[ I_i^{(n)} = 1 \right]. \quad (2.71)$$

We have

$$\begin{aligned} \tau(n) &= \tau_c(n) + \tau_d(n) \\ &= \left\lceil \frac{n\mu}{\bar{W}} \right\rceil + \left\lceil \frac{n\mathcal{E}_D(\tilde{\lambda}) + n^\alpha}{\bar{W}} \right\rceil \\ &< \frac{n}{\bar{W}} (\min\{\lambda^* - \bar{W}, \beta\} + \mathcal{E}_D(\tilde{\lambda}) + n^\alpha/n) + 2 \\ &= \frac{n}{\bar{W}} (\tilde{\lambda} - \bar{W} + \mathcal{E}_D(\tilde{\lambda}) + n^\alpha/n) + 2. \end{aligned} \quad (2.72)$$

Recalling that  $\tilde{\mathcal{E}} = \tilde{\lambda} + \mathcal{E}_D(\tilde{\lambda})$ ,

$$\lim_{n \rightarrow \infty} \frac{n}{n + \tau(n)} \geq \frac{\bar{W}}{\tilde{\lambda} + \mathcal{E}_D(\tilde{\lambda})} = \frac{\bar{W}}{\tilde{\mathcal{E}}}. \quad (2.73)$$

The communication rate would be

$$\rho = \lim_{n \rightarrow \infty} \rho_n = \lim_{n \rightarrow \infty} \frac{nR \mathbb{P} \left[ I_i^{(n)} = 1 \right]}{n + \tau(n)} \quad (2.74a)$$

$$= R \left( \lim_{n \rightarrow \infty} \frac{n}{n + \tau(n)} \right) \left( \lim_{n \rightarrow \infty} \mathbb{P} \left[ I_i^{(n)} = 1 \right] \right) \quad (2.74b)$$

$$\geq \frac{R\bar{W}}{\tilde{\mathcal{E}}}, \quad wp1. \quad (2.74c)$$

Note that the first parentheses in (2.74b) is lower bounded using (2.73) and the second parentheses  $\lim_{n \rightarrow \infty} \mathbb{P} \left[ I_i^{(n)} = 1 \right] = 1$  using (2.66c).

### 2.3.2 High Harvesting Rate

When  $\lambda^* \leq \bar{W}$ , no pre-charging period is required. As the sampling rate is small compared to the arrival energy, the sampling is finished before the packet transmission is over. Note that since  $\lambda^* \leq \bar{W}$ , then  $\lambda^* \leq \bar{W} + \beta$ , so  $\tilde{\lambda} = \lambda^*$ . It will be shown that if sampling is done by time  $\tau_s = \lceil n\lambda^*/\bar{W} \rceil$ , the sampling rate will converge to  $\lambda^*$  in probability.

**Lemma 2.5** *If  $\lambda^* \leq \bar{W}$ , and the sampling is stopped at time  $\tau_s$ ,*

$$\lim_{n \rightarrow \infty} \mathbb{P} [S/n < \lambda^* - n^{\alpha-1}] = 0.$$

**Proof:** For the number of samples taken from a packet we have

$$S \geq \sum_{t=1}^{\tau_s-1} W_t - \sum_{i=1}^{K_s} V_i - U_{\tau_s}, \quad (2.75)$$

where  $K_s$  is the number of busy periods in  $[1, \tau_s]$  and  $V_i$  is the energy discarded in busy period  $i$  due to limited battery.  $U_{\tau_s}$  is the stored energy in the battery at time  $\tau_s$ . Then,

$$\begin{aligned} \mathbb{P} [S/n < \lambda^* - n^{\alpha-1}] &\leq \mathbb{P} \left[ \sum_{t=1}^{\tau_s-1} W_t - \sum_{i=1}^{K_s} V_i - U_{\tau_s} < n\lambda^* - n^\alpha \right] \\ &\leq \mathbb{P} \left[ \sum_{t=1}^{\tau_s-1} W_t - \sum_{i=1}^{\tau_s} V_i - U_{\tau_s} < \left( \left\lceil \frac{n\lambda^*}{\bar{W}} \right\rceil - 1 \right) \bar{W} + \bar{W} - n^\alpha \right], \end{aligned} \quad (2.76)$$

where  $n\lambda^*$  is upper bounded by  $\bar{W} \lceil \frac{n\lambda^*}{\bar{W}} \rceil$ . Also, we used the fact that  $\tau_s \geq K_s$ . Similar to what we did in the proof of Theorem 2.2, we define the random variables  $A_i$ ,  $i = 1, \dots, \tau_s + 2$ , as

$$A_i = -V_i, \quad i = 1, 2, \dots, \tau_s \quad (2.77a)$$

$$A_{\tau_s+1} = -U_{\tau_s} + n^\alpha/2 \quad (2.77b)$$

$$A_{\tau_s+2} = \sum_{t=1}^{\tau_s-1} W_t - (\tau_s - 1)\bar{W} + n^\alpha/2 - \bar{W}. \quad (2.77c)$$



Following the same steps as in (2.60) through (2.61),

$$\mathbb{P}[S/n < \lambda^* - n^{\alpha-1}] \leq \left(\frac{n\lambda^*}{\bar{W}} + 1\right) \exp(-n\beta r^*) + \frac{n^{-\alpha}\eta_W}{2} + \exp\left(-\frac{2(n^\alpha/2 - \bar{W})^2}{(n\lambda^*/\bar{W})b^2}\right),$$

which goes to zero as  $n \rightarrow \infty$ .

□

Decoding is started at time  $\tau_s + 1$  for  $\tau_d$  time units as defined in (2.23). If  $\bar{W} \leq \lambda^* + \mathcal{E}_D(\lambda^*)$ , then

$$\begin{aligned} n &\leq \frac{n\lambda^*}{\bar{W}} + \frac{n\mathcal{E}_D(\lambda^*)}{\bar{W}} \\ &\leq \tau_s + \tau_d. \end{aligned} \tag{2.78}$$

Thus,  $\tau_s + \tau_d$  time units are spent for sampling and decoding. This includes a time interval of length  $\tau_s + \tau_d - n$  following the packet transmission in which decoding is completed. On the other hand, if  $\bar{W} > \lambda^* + \mathcal{E}_D(\lambda^*) + n^{\alpha-1} + 2/n$ , then  $n > \tau_s + \tau_d$ , so not only the pre-charging period is zero, but also no extra time is spent for decoding so the rate would be  $R$ . Note that if for some  $n$ ,  $\lambda^* + \mathcal{E}_D(\lambda^*) < \bar{W} \leq \lambda^* + \mathcal{E}_D(\lambda^*) + n^{\alpha-1} + 2/n$ , then there exists  $n_1$  such that for all  $n > n_1$ ,  $\bar{W} > \lambda^* + \mathcal{E}_D(\lambda^*) + n^{\alpha-1} + 2/n$ . It follows that the throughput is

$$\begin{aligned} \rho &\geq \lim_{n \rightarrow \infty} \frac{nR}{n + [\tau_s + \tau_d - n]^+} \\ &= R \min\{1, \frac{\bar{W}}{\mathcal{E}_*}\}. \end{aligned} \tag{2.79}$$

## Summary of Algorithm

In summary, the receiver calculates the required gap between the packets,  $\tau(n)$  and reports it to the transmitter. This value is calculated using  $\bar{W}$ ,  $\beta$  and the decoding function  $\mathcal{E}_D(\lambda)$  which yields  $\lambda^*$  and  $\mathcal{E}_D(\lambda^*)$ . The proposed algorithm is as follows:

- If  $\lambda^* > \bar{W}$ , the transmitter sends the new packets every  $\tau_c + n + \tau_d$  time units.

The receiver charges the battery during  $\tau_c$ , samples the packet with sampling rate  $\tilde{\lambda}$  and decodes the packet during  $\tau_d$ .

- If  $\lambda^* \leq \bar{W}$ , the transmitter sends the new packets every  $n + [\tau_s + \tau_d - n]^+$  time units. Starting from zero energy in the battery, the receiver takes samples until time  $\tau_s$  and then starts decoding.

**Theorem 2.3** *In a variable-timing transmission system with packets encoded at rate  $R$  while the energy arrival at the receiver is i.i.d. with mean  $\bar{W}$  and the battery capacity is  $n\beta$  for an  $n$ -length code, the communication rate of*

$$\rho = R \min\left\{1, \frac{\bar{W}}{\tilde{\mathcal{E}}}\right\}$$

*is achievable where  $\tilde{\lambda} = \min\{\lambda^*, \beta + \bar{W}\}$  and  $\tilde{\mathcal{E}} = \tilde{\lambda} + \mathcal{E}_D(\tilde{\lambda})$ .*

## 2.4 Outerbound

We would like to show that under any policy the achievable rate in Theorem 2.3 is optimum for large number of packets and large number of symbols in a packet. It is obvious that the communication rate never exceeds  $R$ . So, we need to prove that it doesn't exceed  $R\bar{W}/\tilde{\mathcal{E}}$ .

We start by defining

$$\hat{\lambda}_n = \min\{\lambda^*, \beta + \bar{W} + n^{\alpha-1}\} \quad (2.80)$$

and

$$\hat{\mathcal{E}}_n = \hat{\lambda}_n + \mathcal{E}_D(\hat{\lambda}_n) \quad (2.81)$$

as a function of the block length  $n$ . Note that  $\lim_{n \rightarrow \infty} \hat{\mathcal{E}}_n = \tilde{\mathcal{E}}$ . We may drop  $n$  when the dependency on  $n$  is clear from the text.

Let's consider an arbitrary policy in which  $M_n$  packets are decoded in the time interval  $[1, T]$  obtaining the communication rate

$$\rho_n(T) = \frac{nRM_n}{T}. \quad (2.82)$$

We prove the following outerbound.

**Theorem 2.4**

$$\mathbb{P} \left[ \rho_n(T) \geq \frac{R\bar{W}}{\tilde{\mathcal{E}}} + \delta(n, T) \right] \leq \epsilon_A(n, T) + \epsilon_2(T),$$

where

$$\delta(n, T) = \frac{R(\bar{W} + 2T^{\alpha-1})}{(1 - \epsilon_1(n))\hat{\mathcal{E}}} - \frac{R\bar{W}}{\tilde{\mathcal{E}}},$$

$$\epsilon_1(n) = \exp \left( -\frac{2n^{2\alpha-1}}{b^2} \right),$$

$$\epsilon_2(T) = \exp \left( -\frac{2T^{2\alpha-1}}{b^2} \right),$$

$$\epsilon_A(n, T) = \exp \left( -\frac{2(1 - \epsilon_1(n)) \left( T^\alpha - n(1 - \epsilon_1(n))\hat{\mathcal{E}} \right)^2}{(T(\bar{W} + 2T^{\alpha-1})) (n\hat{\mathcal{E}})} \right).$$

Note that for  $0.5 < \alpha < 1$ ,  $\lim_{n \rightarrow \infty} \lim_{T \rightarrow \infty} \delta(n, T) = 0$  and  $\lim_{n \rightarrow \infty} \lim_{T \rightarrow \infty} \epsilon_A(n, T) = 0$ . Also,  $\lim_{n \rightarrow \infty} \epsilon_1(n) = 0$  and  $\lim_{T \rightarrow \infty} \epsilon_2(T) = 0$ .

Defining

$$A_n = \left\lfloor \frac{T(\bar{W} + 2T^{\alpha-1})}{\hat{n}_n \hat{\mathcal{E}}} \right\rfloor, \tag{2.83}$$

where  $\hat{n}_n = n(1 - \epsilon_1(n))$ , we have

$$\begin{aligned} \frac{nRA_n}{T} &= \frac{nR}{T} \left\lfloor \frac{T(\bar{W} + 2T^{\alpha-1})}{\hat{n}_n \hat{\mathcal{E}}} \right\rfloor \\ &\leq \frac{R(\bar{W} + 2T^{\alpha-1})}{(1 - \epsilon_1(n))\hat{\mathcal{E}}} \end{aligned} \tag{2.84}$$

$$= \frac{R\bar{W}}{\tilde{\mathcal{E}}} + \delta(n, T), \tag{2.85}$$

where  $R\bar{W}/\tilde{\mathcal{E}}$  is added and subtracted in going from (2.84) to (2.85). It then follows from (2.82) that

$$\mathbb{P} \left[ \rho_n(T) \geq \frac{R\bar{W}}{\tilde{\mathcal{E}}} + \delta(n, T) \right] \leq \mathbb{P} [M_n \geq A_n]. \tag{2.86}$$

So, to prove the claim, it is enough to show that

$$\mathbb{P}[M_n \geq A_n] \leq \epsilon_A(n, T) + \epsilon_2(T). \quad (2.87)$$

**Proof:** Denoting the arrival energy in the transmission block of decoded packet  $i$  as  $Q_i$ , (2.6) implies that

$$\mathbb{P}[Q_i \geq n\bar{W} + n^\alpha] \leq \exp\left(-\frac{2n^{2\alpha-1}}{b^2}\right) = \epsilon_1(n), \quad (2.88)$$

where  $\alpha \in (0.5, 1)$ . We define

$$X_i = n\hat{\mathcal{E}} \mathbf{1}(Q_i < n\bar{W} + n^\alpha). \quad (2.89)$$

Thus (2.88) implies

$$\mathbb{P}[X_i = 0] \leq \epsilon_1(n), \quad (2.90)$$

Note that  $Q_i$  and consequently  $X_i$  are i.i.d. random variables. Note that  $X_i$  is a lower bound for the energy expended in sampling and decoding packet  $i$ . To see this, suppose energy expended for packet  $i$  were  $X_i$ . In the atypical case, when the arrival energy is plentiful and  $Q_i \geq n\bar{W} + n^\alpha$ , the energy cost  $X_i$  of sampling and decoding is taken to be zero. On the other hand,  $X_i = n\hat{\mathcal{E}}$  corresponds to the typical case when  $Q_i < n\bar{W} + n^\alpha$ . In this case, the available energy for sampling cannot exceed  $n\beta + n\bar{W} + n^\alpha$ . As a result, the sampling rate is limited as

$$\lambda < \beta + \bar{W} + n^{\alpha-1} = \hat{\lambda}_n. \quad (2.91)$$

As depicted in Fig. 2.4, the total energy is convex and non-increasing in  $\lambda$  for  $\lambda < \lambda^*$ . Thus, the expended energy cannot be below  $n\hat{\mathcal{E}}$ . So,  $X_i = n\hat{\mathcal{E}}$  is a lower bound for the expended energy for this case.

We note that to have  $M_n \geq A_n$ , it is necessary to have at least  $A_n$  packets decoded.

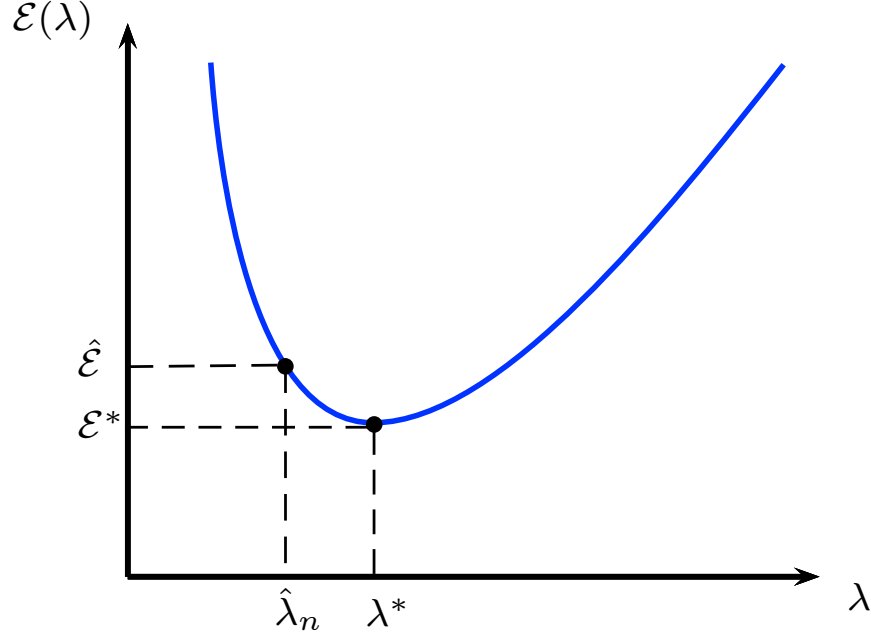


Figure 2.4: Total energy as a function of the sampling rate when  $\hat{\lambda}_n = \beta + \bar{W} + n^{\alpha-1} < \lambda^*$ .

Hence, energy conservation dictates that

$$\mathbb{P}[M_n \geq A_n] \leq \mathbb{P}\left[\sum_{i=1}^{A_n} X_i \leq \sum_{t=1}^T W_t\right]. \quad (2.92)$$

We define the event

$$D = \left\{ \sum_{t=1}^T W_t \leq T\bar{W} + T^\alpha \right\}, \quad (2.93)$$

to represent a typical energy arrival over the total  $T$  time slots. Applying the law of total probability to (2.93) yields

$$\begin{aligned} \mathbb{P}\left[\sum_{i=1}^{A_n} X_i \leq \sum_{t=1}^T W_t\right] &= \mathbb{P}\left[\left\{\sum_{i=1}^{A_n} X_i \leq \sum_{t=1}^T W_t\right\} \cap D\right] + \mathbb{P}\left[\left\{\sum_{i=1}^{A_n} X_i \leq \sum_{t=1}^T W_t\right\} \cap D^c\right] \\ &\leq \mathbb{P}\left[\sum_{i=1}^{A_n} X_i \leq T\bar{W} + T^\alpha\right] + \mathbb{P}[D^c]. \end{aligned} \quad (2.94)$$

According to the Hoeffding's inequality (2.6), we have

$$\mathbb{P}[D^c] \leq \exp\left(-\frac{2T^{2\alpha-1}}{b^2}\right) = \epsilon_2(T). \quad (2.95)$$

Note that  $\mathbb{P}[X_i = 0] = \epsilon_0 \leq \epsilon_1(n)$  and  $\mathbb{P}[X_i = n\hat{\mathcal{E}}] = 1 - \epsilon_0 > 1 - \epsilon_1(n)$ . We lower bound the energy expenditure of the receiver by a fictitious enhanced system that expends energy  $Y_i$  to sample and decode the decoded packet  $i$ . If  $X_i = 0$ , then,  $Y_i = X_i$ . On the other hand, if  $X_i = n\hat{\mathcal{E}}$ ,  $Y_i = n\hat{\mathcal{E}}$  with probability  $p$  and  $Y_i = 0$  with probability  $1 - p$ , where  $p = (1 - \epsilon_1(n))/(1 - \epsilon_0)$ . Thus,  $Y_i$  has PMF

$$\mathbb{P}[Y_i = y] = \begin{cases} \epsilon_1(n) & y = 0, \\ 1 - \epsilon_1(n) & y = n\hat{\mathcal{E}}. \end{cases} \quad (2.96)$$

This implies

$$\mathbb{P}\left[\sum_{i=1}^{A_n} X_i \leq T\bar{W} + T^\alpha\right] \leq \mathbb{P}[C_n], \quad (2.97)$$

where

$$\mathbb{P}[C_n] = \mathbb{P}\left[\sum_{i=1}^{A_n} Y_i \leq T\bar{W} + T^\alpha\right]. \quad (2.98)$$

Then, noting that  $Y_i$ 's are i.i.d., and having

$$\mathbb{E}\left[\sum_{i=1}^{A_n} Y_i\right] = A_n \mathbb{E}[Y_i] = A_n n (1 - \epsilon_1(n)) \hat{\mathcal{E}} = A_n \hat{n}_n \hat{\mathcal{E}}, \quad (2.99)$$

we would like to upper bound (2.97) using the Hoeffding inequality (2.6). So, we rewrite (2.97) as

$$\mathbb{P}[C_n] = \mathbb{P}\left[\sum_{i=1}^{A_n} Y_i \leq A_n \hat{n}_n \hat{\mathcal{E}} + \left(T\bar{W} + T^\alpha - A_n \hat{n}_n \hat{\mathcal{E}}\right)\right] \quad (2.100)$$

By (2.83),  $A_n > \frac{T(\bar{W} + 2T^{\alpha-1})}{\hat{n}_n \hat{\mathcal{E}}} - 1$ , so,

$$\begin{aligned} \mathbb{P}[C_n] &\leq \mathbb{P}\left[\sum_{i=1}^{A_n} Y_i \leq A_n \hat{n}_n \hat{\mathcal{E}} + \left(T\bar{W} + T^\alpha - \left(\frac{T(\bar{W} + 2T^{\alpha-1})}{\hat{n}_n \hat{\mathcal{E}}} - 1\right) \hat{n}_n \hat{\mathcal{E}}\right)\right] \\ &= \mathbb{P}\left[\sum_{i=1}^{A_n} Y_i \leq A_n \hat{n}_n \hat{\mathcal{E}} - \left(T^\alpha - \hat{n}_n \hat{\mathcal{E}}\right)\right] \end{aligned} \quad (2.101)$$

Now, we apply the Hoeffding inequality (2.6) to (2.101), yielding

$$\begin{aligned} \mathbb{P}[C_n] &\leq \exp\left(-\frac{2(T^\alpha - \hat{n}_n \hat{\mathcal{E}})^2}{A_n (n\hat{\mathcal{E}})^2}\right) \\ &\leq \exp\left(-\frac{2(1 - \epsilon_1(n))(T^\alpha - \hat{n}_n \hat{\mathcal{E}})^2}{T(\bar{W} + 2T^{\alpha-1})(n\hat{\mathcal{E}})}\right) = \epsilon_A(n, T). \end{aligned} \quad (2.102)$$

where in (2.102), we used the fact that (2.83) implies  $A_n \leq \frac{T(\bar{W} + 2T^{\alpha-1})}{\hat{n}_n \hat{\mathcal{E}}}$ .  $\square$

We conclude that

$$\lim_{n \rightarrow \infty} \lim_{T \rightarrow \infty} \mathbb{P}\left[\rho_n(T) \leq \frac{R\bar{W}}{\hat{\mathcal{E}}} + \delta(n, T)\right] = 0, \quad (2.103)$$

which proves the optimality of the achievable schemes.

## 2.5 Optimum Code Rate

So far we studied the maximum communication rate for a fixed code rate  $R$  with variable-timing transmission. Now, we look at the maximization problem over both sampling rate and the code rate. Following the model in [27,28], we assume the decoding energy is a function,  $f$  of  $R/C\lambda$ , that is  $f(R/C\lambda) = \mathcal{E}_D(\lambda)$  for a fixed  $R$  and  $C$ . According to Theorem 2.3, for a fixed  $R$  and  $\lambda$ ,  $C$  and  $\bar{W}$ , the rate of  $\bar{W}R/(\lambda + f(R/C\lambda))$

is achievable. We wish to maximize the achievable rate as follows.

$$\check{\rho} = \max_{R, \lambda} \frac{\bar{W}R}{\lambda + f(R/C\lambda)} \quad (2.104a)$$

$$\text{s.t. } 0 < R < C\lambda \quad (2.104b)$$

$$\lambda \leq \lambda_{\max}, \quad (2.104c)$$

where

$$\lambda_{\max} \triangleq \min \{ \beta + \bar{W}, 1 \}. \quad (2.105)$$

We assume that  $\mathcal{E}_D(R/C\lambda)$  is zero only at  $R = 0$ . Also we assume that the function of  $f(z)$ , where  $z = R/C\lambda$ , is differentiable.

According to the KKT optimality conditions, complementary slackness implies that the Lagrange multipliers corresponding to the strict inequalities should be zero. Defining  $\mu$  as the Lagrange multiplier, the Lagrangian is

$$L(\lambda, R, \mu) = \frac{\bar{W}R}{f(R/C\lambda) + \lambda} - \mu(\lambda_{\max} - \lambda). \quad (2.106)$$

If  $\mu = 0$ , then,  $\partial L / \partial \lambda = 0$  and  $\partial L / \partial R = 0$ , yielding

$$1 - \frac{R}{C\lambda^2} f' \left( \frac{R}{C\lambda} \right) = 0 \quad (2.107a)$$

$$f \left( \frac{R}{C\lambda} \right) + \lambda - \frac{R}{C\lambda} f' \left( \frac{R}{C\lambda} \right) = 0. \quad (2.107b)$$

Eq. (2.107) gives  $f(R/C\lambda) = 0$ , implying  $\check{R} = 0$ , which conflicts with our modeling assumptions. So,  $\mu \neq 0$  and, according to complementary slackness, the constraint (2.104c) should be active at the optimum point. So,  $\check{\lambda} = \lambda_{\max}$  at this point and

$$\left. \frac{d}{dR} \left( \frac{\bar{W}R}{f(R/C\lambda_{\max}) + \lambda_{\max}} \right) \right|_{R=\check{R}} = 0. \quad (2.108)$$



This implies

$$f\left(\frac{\check{R}}{C\lambda}\right) = f'\left(\frac{\check{R}}{C\lambda}\right) \frac{\check{R}}{C\lambda_{\max}} - \lambda_{\max}. \quad (2.109)$$

Then, for the optimum communication rate we will have

$$\check{\rho} = \frac{\bar{W}C\lambda_{\max}}{f'(\check{R}/C\lambda_{\max})}.$$

Fig. 2.5 shows how the communication rate changes with the code rate. Here, we used the decoding energy model in [27], that is  $f(z) = \frac{1}{1-z} \log\left(\frac{1}{1-z}\right)$ . At each code rate, the total energy is minimized. As the code rate goes up, the optimum sampling also goes up until some point where it reaches  $\beta + \bar{W}$  which is the maximum achievable sampling rate. From this point on, the sampling rate is fixed at maximum and increasing  $R$ , shrinks the capacity gap and increases the decoding energy. For a short while, increasing  $R$  would be dominant and although the decoding energy is growing, still the communication rate goes up but after reaching the peak, the growth of the decoding energy would dominate and the communication rate falls down. We can see that the total energy is a non-decreasing convex function of the code rate.

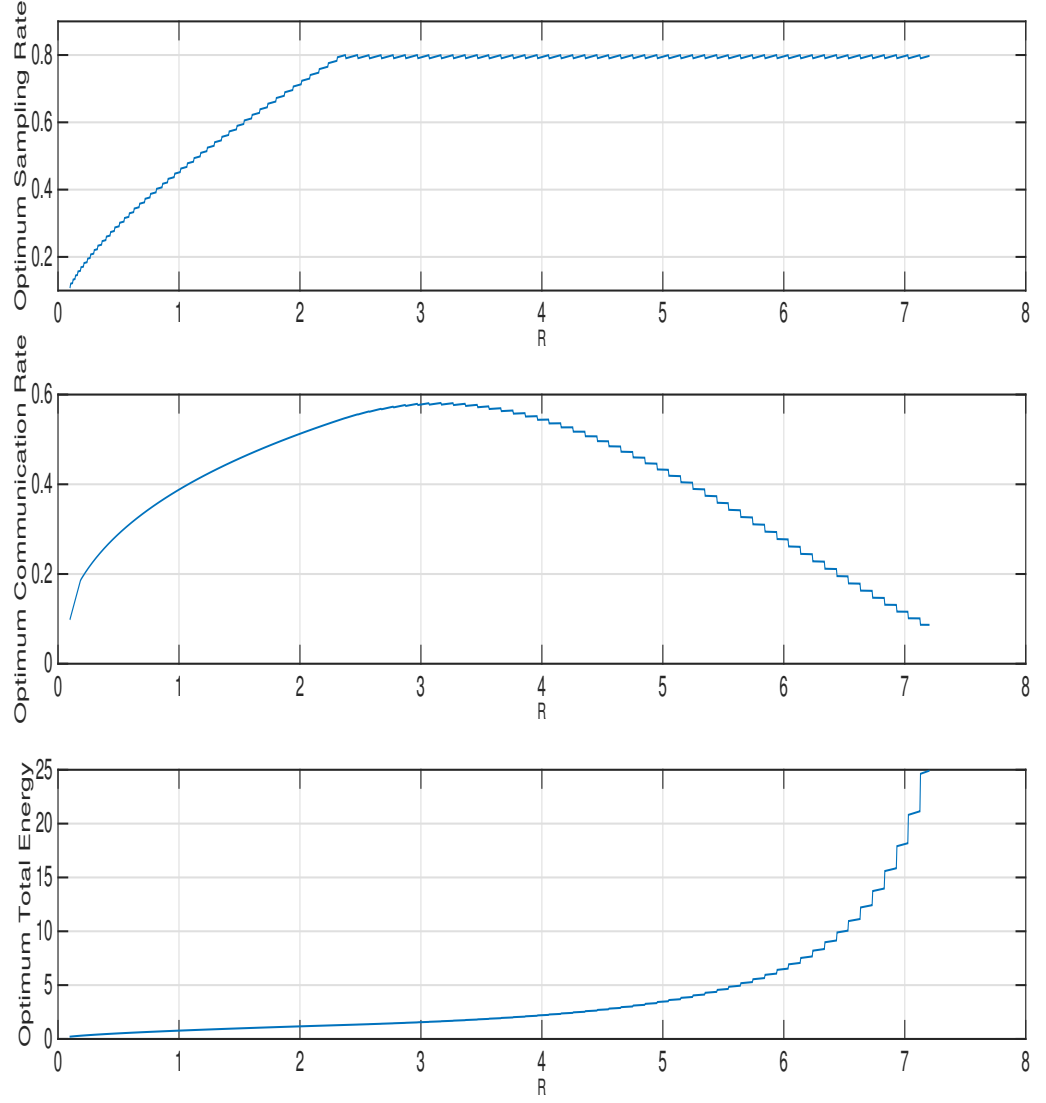


Figure 2.5: Sampling rate, communicate rate and the total normalized energy versus  $R$  (bits/s/Hz) for  $\bar{W} = 0.3\mu J$ ,  $\beta = 0.5\mu J$ ,  $C = 10$  bits/s/Hz,  $n = 10^6$ .

## Chapter 3

### Energy-Harvesting Receivers in Fading Channels

We consider the operation of an energy harvesting receiver in a point-to-point fading channel with additive white Gaussian noise (AWGN). Knowledge of the channel state information is not available at the transmitter. The limited rate of harvesting energy at the receiver along with the time variation of the channel can degrade the performance of the system. However, we will show that channel state knowledge at the receiver can improve the performance of the system. We propose a channel-selective sampling strategy that optimizes a tradeoff between the energy costs of sampling and decoding at the receiver. Based on this tradeoff, we derive a policy maximizing the communication rate and we characterize an energy-constrained rate region. We extend the results to the receivers with finite battery capacity. We will consider the fading channel under two models: fast-fading and slow-fading. In the fast-fading case the channel gain from symbol to symbol period is an i.i.d. sequence; but in the slow-fading case, the channel is fixed during a block length and it changes i.i.d. to the next block.

#### 3.1 System Model

In this chapter, we consider a fast fading AWGN channel such that the channel gain,  $H_t$ , is fixed during one symbol period but changes independently and identically in the next symbol period. The relationship between the transmit and received signal at symbol period  $t$  is

$$Y_t = H_t X_t + Z_t, \quad (3.1)$$

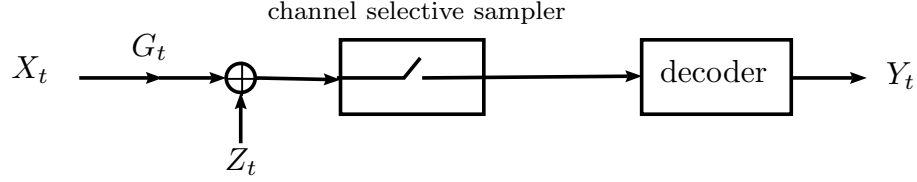


Figure 3.1: Channel Model

where  $Z_t$  is the AWGN noise at time  $t$ . We define

$$G_t = \|H_t\|^2 \quad (3.2)$$

as the instantaneous channel gain.

We consider an idealized model in which the estimate of the channel  $H_t$  is available before the arrival of the sample  $x_t$ . In such an idealized setting, acquiring channel state information at the receiver doesn't affect the sampling rate and the final communication rate. We also assume that the block length,  $n$ , is long enough to realize the ergodic variation of the fading channel.

The capacity of this physical channel depends on the performance of the receiver front-end, which is coupled to the sampling energy. In this system, it is useful to think of the sampling energy per sampled symbol  $\nu$  fixed for a physical channel of capacity  $C$ . Without loss of generality, we assume  $\nu = 1$ .

According to [47], the ergodic capacity of the fading channel (3.1) with average transmitted power  $p = \mathbb{E}[\|X_t\|^2]$  is

$$C = \mathbb{E} \left[ \log \left( 1 + \frac{pG_t}{\sigma^2} \right) \right]. \quad (3.3)$$

Thus, the gap to the capacity would be

$$\delta = 1 - \frac{R}{\lambda \mathbb{E} \left[ \log \left( 1 + \frac{pG_t}{\sigma^2} \right) \right]}. \quad (3.4)$$

This model extends the model in Chapter 2 where the channel is time-invariant. Here, we propose a channel selective sampling strategy which dictates correlation between

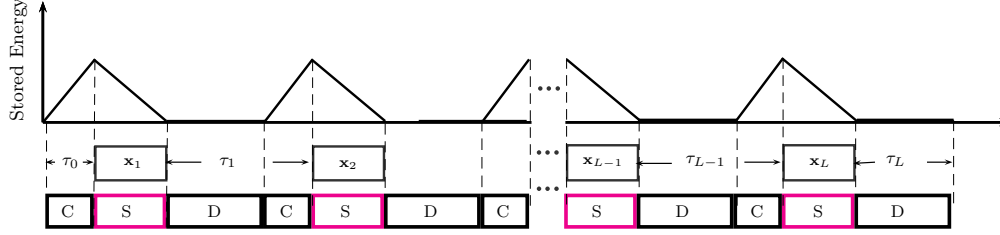


Figure 3.2: Variable-timing optimum policy: Transmitted packets are labeled  $\mathbf{x}_1, \mathbf{x}_2, \dots$  while intervals marked “C,” “S,” and “D” mark when the receiver is (C)harging the battery, (S)ampling a packet, and (D)ecoding that packet. The corresponding graph depicts the receiver’s stored energy.

the sampling rate and the channel gain and based on this model we will optimize the communication rate.

### 3.1.1 Energy Harvesting Model

Here we assume that energy  $\bar{W}$  arrives deterministically in every symbol period. We believe this is an appropriate model when code words are transmitted in milliseconds and the coherence time of the energy harvesting process is on the order of minutes or hours. On the other hand, when codewords are much longer than the harvesting coherence time, it also can be shown that system performance chiefly depends on the average harvesting rate  $\bar{W} = \mathbb{E}[W_t]$ .

## 3.2 Achievability: Channel Selective Sampling

In the variable-timing transmission setting, the packets are transmitted apart by some time  $\tau$  which is reserved for decoding the previous sampled packet and also collecting energy for sampling the current packet; see Fig. 3.2. What we need to characterize for such a protocol is the time interval between two transmissions,  $\tau$  which is computed at the receiver and fed to the transmitter just one time before starting the communication.

Based on our discussion in previous chapters, we know that not only sampling a subset of the symbols is enough but also it may save energy at the receiver, leading to a better throughput in energy-limited receivers. Furthermore, to take advantage of receiver CSI in the sampling process in order to enhance the performance of the system,

we may wish to sample a subset of the symbols for which the channel is good. However, how do we know if there will be enough energy available to take those samples? This is more apparent when the process of harvesting energy is stochastic. However, even in a deterministic energy arrival setting when the rate of harvesting energy is limited, energy may not be available to take samples when the channel is good.

Another point to be addressed is that how large should the channel gain be so that the corresponding sample is taken? Assume we set a threshold,  $\gamma$ , on the channel gain. Any symbol which experiences a channel gain above this threshold is taken and the rest are dropped. We would like to find the optimum threshold,  $\gamma^*$  which maximizes the rate,  $\rho$ .

We define the new random process  $\hat{G}_t$  as

$$\hat{G}_t(\gamma) \triangleq \begin{cases} G_t & G_t \geq \gamma, \\ 0 & \text{otherwise.} \end{cases} \quad (3.5)$$

Note that  $\hat{G}_t(\gamma)$  is an i.i.d. sequence, with each sample identical to  $\hat{G}(\gamma)$ . Also, note that the number of samples taken through this strategy is a random variable which is a function of the threshold gain,  $\gamma$ .

When the channel (3.1) is concatenated with the channel selective sampler as in Fig. 3.1, the capacity of the effective channel is

$$\hat{C}(\gamma) = \mathbb{E} \left[ \log \left( 1 + \frac{p\hat{G}(\gamma)}{\sigma^2} \right) \right]. \quad (3.6)$$

Under the channel selective sampling, the capacity gap  $\delta$  becomes the function of the sampling threshold,  $\gamma$ . Specifically,

$$\delta(\gamma) = 1 - \frac{R}{\hat{C}(\gamma)}. \quad (3.7)$$

Thus the decoding energy is also a function of the gap to the capacity. In this chapter,

we use  $f(\cdot)$  to describe this relationship. We also assume  $\mathcal{E}_D(\cdot)$  is a function of  $\gamma$ . Then,

$$\mathcal{E}_D(\gamma) = f(\delta(\gamma)) = f\left(1 - R/\hat{C}(\gamma)\right). \quad (3.8)$$

As an example, according to the conjecture in [27],  $f(\delta) = (\alpha/\delta) \log(1/\delta)$ , so

$$\mathcal{E}_D(\gamma) = \frac{\alpha}{\delta(\gamma)} \log \frac{1}{\delta(\gamma)} \quad (3.9)$$

for some coefficient  $\alpha$ . Then, using (3.7),

$$\mathcal{E}_D(\gamma) = \frac{\alpha}{1 - R/\hat{C}(\gamma)} \log \left( \frac{1}{1 - R/\hat{C}(\gamma)} \right). \quad (3.10)$$

Assuming harvested energy is sufficient, the number of taken samples,  $S$ , is

$$S = \sum_{t=1}^n \mathbf{1}(G_t \geq \gamma). \quad (3.11)$$

According to the law of large numbers,

$$\lim_{n \rightarrow \infty} \frac{S}{n} = \mathbb{E}[\mathbf{1}(G_t \geq \gamma)] = \Pr[G \geq \gamma] \quad \text{w.p.1.} \quad (3.12)$$

Thus, the sampling rate, or the normalized sampling energy is equal to  $\lambda = \Pr[G \geq \gamma]$ . As the threshold level  $\gamma$  is increased, fewer samples are taken, so the capacity (3.6) is decreased. This relationship is depicted in Fig. 3.3 where we have used Rayleigh fading model. The channel gain is unit-mean exponential and the noise zero-mean and unit variance. In this figure, the transmit power is unity and  $R = 0.4$ .

It can be seen that increasing  $\gamma$ , will decrease the capacity, shrinking the gap between rate and the capacity which consequently leads to an increase in the decoding energy. While increasing  $\gamma$  will increase the decoding energy, it makes it less likely that the gain of the channel be over the threshold which reduces the sampling energy. Thus,  $\Pr[G \geq \gamma]$  decreases. So, with changing  $\gamma$  the tradeoff curve in Fig. 3.4 is obtained. We have used the model in [27] for LDPC decoder complexity as in (3.9) and (3.10).

Due to the convexity of the decoding energy, the total energy in Fig. 3.4, the total

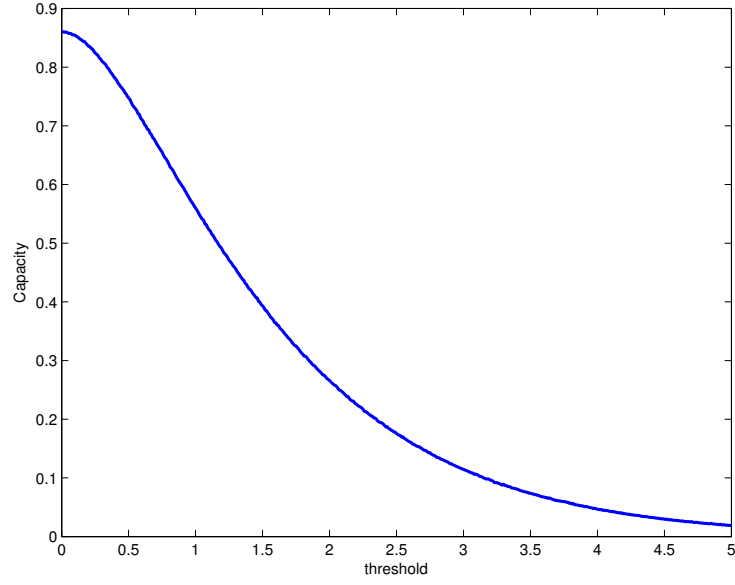


Figure 3.3: Capacity of the effective channel under channel selective sampling versus the threshold gain in Rayleigh fading.

energy requirement is minimum at the point where the 45-degree line is tangent to the decoding energy curve. We call the sampling rate at this point  $\lambda^*$ . We wish to select the channel gain threshold such that the sampling rate is  $\lambda^*$ . So, we choose  $\gamma^*$  such that

$$\lambda^* = \Pr[G \geq \gamma^*] \quad (3.13)$$

or in other words

$$\gamma^* = \arg \min_{\gamma} \Pr[G \geq \gamma] + \mathcal{E}_D(\gamma), \quad (3.14)$$

Setting the corresponding threshold to  $\gamma^*$ , the decoding energy is  $\mathcal{E}_D(\gamma^*)$  and, the total minimum energy is

$$\begin{aligned} \mathcal{E}^* &= \Pr[G \geq \gamma^*] + \mathcal{E}_D(\gamma^*) \\ &= \lambda^* + \mathcal{E}_D^*, \end{aligned} \quad (3.15)$$



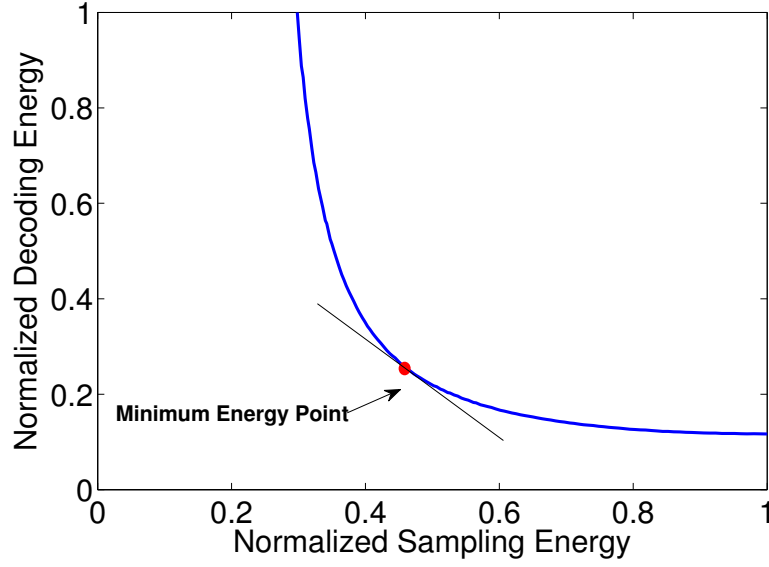


Figure 3.4: Normalized decoding energy,  $\mathcal{E}_D(\gamma)$ , versus sampling energy,  $\mathcal{E}_S(\gamma) = \Pr[G \geq \gamma]$  for a fixed  $R$  under channel selective sampling strategy. The total energy requirement is minimum at Minimum Energy Point.

where  $\mathcal{E}_D^* = \mathcal{E}_D(\gamma^*)$ .

We will show that with appropriate choice of  $\tau$ , we can make sure that enough energy is available to achieve the sampling rate  $\lambda^*$ , under the channel selective sampling policy with the parameter threshold  $\gamma^*$ .

### 3.3 Variable Timing: Achievable Rates

To be able to sample at  $\lambda^*$ , we need the stored battery energy and energy harvesting rate to be sufficiently large. If  $\lambda^* > \bar{W}$ , the energy collected during the packet transmission is not enough for sampling; so it is required to charge the battery before the packet transmission. As  $\bar{W}$  units of energy arrives in each symbol period, we require the initial energy at the start of receiving the packet be

$$U_0 = n\lambda^* - n\bar{W}. \quad (3.16)$$

Thus it takes

$$\tau_S = (n\lambda^* - n\bar{W})/\bar{W} \quad (3.17)$$

time units to collect this energy. Also, it takes

$$\tau_D = n\mathcal{E}_D^*/\bar{W} \quad (3.18)$$

time units to collect energy for decoding the previous sampled packet. So,

$$\tau = \tau_S + \tau_D = \frac{(n\lambda^* - n\bar{W}) + n\mathcal{E}_D^*}{\bar{W}} \quad (3.19)$$

and the communication rate is

$$\begin{aligned} \rho &= \frac{nR}{(n\lambda^* - n\bar{W} + n\mathcal{E}_D^*)/\bar{W} + n} \\ &= \frac{R\bar{W}}{\lambda^* + \mathcal{E}_D^*} \\ &= \frac{R\bar{W}}{\mathcal{E}^*}. \end{aligned} \quad (3.20)$$

A timing diagram of this achievable scheme is shown in Fig. 3.2.

On the other hand, if  $\bar{W} \geq \lambda^*$ , it is not necessary to collect energy and even there is extra energy arriving during symbol arrivals which can be expended for decoding. So, while  $\tau_S = 0$ , if  $\mathcal{E}_D^* + \lambda^* > \bar{W}$

$$\tau_D = \frac{n\mathcal{E}_D^* - n\bar{W} + n\lambda^*}{\bar{W}}.$$

time units are required for decoding. Consequently, the rate would be

$$\rho = \frac{nR}{(n\mathcal{E}_D^* - n\bar{W} + n\lambda^*)/\bar{W} + n} = \frac{R\bar{W}}{\mathcal{E}^*}. \quad (3.21)$$

If  $\bar{W} \geq \mathcal{E}_D^* + \lambda^*$ , then no extra time other than the packet arrival time interval,  $n$ , is required and the rate would be  $\rho = R$ .

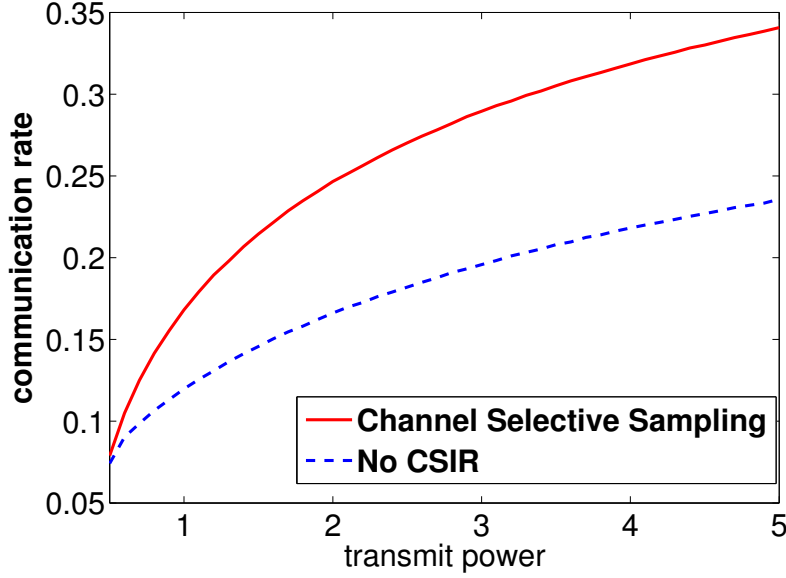


Figure 3.5: The achievable rate for channel selective sampling strategy and the strategy at which the sampling is optimized independently from channel with optimum parameters in terms of the transmit power.

Fig. 3.5, compares the performance of this policy with the one with no receiver CSI. We will show that the channel selective policy with the parameter threshold  $\gamma^*$  is optimum.

### 3.4 Outerbound

Assume an arbitrary policy decoding  $M$  packets in  $T$  time slots. Assume  $\lambda_i$  fraction of symbols are sampled from packet  $i$  while  $\mathcal{E}_i$  units of energy are expended for sampling and decoding it. We define  $a_i = \lfloor n\lambda_i \rfloor$ . Let  $\mathbb{A} \subset \{1, 2, \dots, n\}$  denote the subset of time slots  $t$  with the best  $a_i$  channel gains. We use  $Y_1^{(n)}, Y_2^{(n)}, \dots, Y_n^{(n)}$  to denote the random variables  $G_1, G_2, \dots, G_n$  sorted in decreasing order. We define the threshold  $\gamma(\lambda_i)$  such that  $\mathbb{P}[G_t \geq \gamma(\lambda_i)] = \lambda_i$ . We start by verifying the following lemma.

**Lemma 3.1**  $Y_{a_i}^{(n)}$  converges in distribution to  $\gamma(\lambda_i)$  as  $n \rightarrow \infty$  where  $\mathbb{P}[G_t \geq \gamma(\lambda_i)] = \lambda_i$ .

**Proof:** Let  $N(y) = |\{t \in [1, n] | G_t \geq y\}|$  denote the number of samples of  $G_t$  greater than  $y$ . Since the  $G_t$  are i.i.d.,  $N(y)$  has the binomial distribution

$$\mathbb{P}[N(y) = j] = \binom{n}{j} \bar{F}_G(y)^j F_G(y)^{n-j}, \quad (3.22)$$

where  $F_G(\cdot)$  and  $\bar{F}_G(y)$  denote CDF and complementary CDF, respectively. Note that  $N(y)$  has expected value and variance

$$\mu_n = n\bar{F}_G(y), \quad \sigma_n^2 = nF_G(y)\bar{F}_G(y). \quad (3.23)$$

Now we observe that

$$\mathbb{P}[Y_{a_i}^{(n)} \geq y] = \mathbb{P}[N(y) \geq a_i] = \mathbb{P}\left[\frac{N(y) - \mu_n}{\sigma_n} \geq \frac{a_i - \mu_n}{\sigma_n}\right].$$

As  $n \rightarrow \infty$ , the binomial CDF of  $N(y)$  approaches a Gaussian  $(\mu_n, \sigma_n^2)$  distribution.

Thus

$$\lim_{n \rightarrow \infty} \mathbb{P}[Y_{a_i}^{(n)} \geq y] = \lim_{n \rightarrow \infty} \mathcal{Q}\left[\frac{a_i - n\bar{F}_G(y)}{\sqrt{nF_G(y)\bar{F}_G(y)}}\right]. \quad (3.24)$$

Since  $\lim_{n \rightarrow \infty} a_i/n = \lambda_i$ ,

$$\lim_{n \rightarrow \infty} \mathbb{P}[Y_{a_i}^{(n)} \geq y] = \lim_{n \rightarrow \infty} \mathcal{Q}\left[\frac{n(\lambda_i - \bar{F}_G(y))}{\sqrt{nF_G(y)\bar{F}_G(y)}}\right]. \quad (3.25)$$

If  $y > \gamma(\lambda_i)$ , then  $\bar{F}_G(y) < \lambda_i$ . This implies  $\lim_{n \rightarrow \infty} \mathbb{P}[Y_{a_i}^{(n)} \geq y] = 0$ . On the other hand, if  $y < \gamma(\lambda_i)$ , then  $\bar{F}_G(y) > \lambda_i$  and  $\lim_{n \rightarrow \infty} \mathbb{P}[Y_{a_i}^{(n)} \geq y] = 1$ .  $\square$

We define  $I_i(G_t)$  as the mutual information of symbol  $t$  in the decoded packet  $i$ . The total normalized energy expended for packet  $i$ ,  $\mathcal{E}_i$ , can be lower bounded as

$$\mathcal{E}_i \geq \lambda_i + h\left(\frac{\sum_{t \in \mathbb{A}} I_i(G_t)}{n}\right), \quad (3.26)$$

where  $h(\cdot)$  is the decoding energy as a function of the accumulated mutual information. In other words,  $h(x) = f(1 - R/x)$ . Noting that if symbol  $t \in \mathbb{A}$ , then  $\mathbf{1}(G_t \geq Y_{a_i}^{(n)}) = 1$ , for some  $\delta > 0$  we have

$$\frac{1}{n} \sum_{t \in \mathbb{A}} I_i(G_t) = \frac{1}{n} \sum_{t=1}^n I_i(G_t) \mathbf{1}(G_t \geq Y_{a_i}^{(n)}) \quad (3.27)$$

$$\begin{aligned} &\leq \frac{1}{n} \sum_{t=1}^n \left[ I_i(G_t) \mathbf{1}(G_t \geq \gamma(\lambda_i) - \delta) + I_i(G_t) \mathbf{1}(Y_{a_i}^{(n)} \leq G_t < \gamma(\lambda_i) - \delta) \right] \\ &\leq \frac{1}{n} \sum_{t=1}^n I_i(G_t) \mathbf{1}(G_t \geq \gamma(\lambda_i) - \delta) + \frac{1}{n} I_i(\gamma(\lambda_i)) \sum_{t=1}^n \mathbf{1}(Y_{a_i}^{(n)} \leq G_t < \gamma(\lambda_i) - \delta), \end{aligned} \quad (3.28)$$

where in the second term,  $I_i(G_t)$  is upper bounded by  $I_i(\gamma(\lambda_i))$ . According to the strong law of large numbers, for the first term in (3.28) we have

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t=1}^n I_i(G_t) \mathbf{1}(G_t \geq \gamma(\lambda_i) - \delta) = \mathbb{E}[I_i(G) | G \geq \gamma(\lambda_i) - \delta] \mathbb{P}[G \geq \gamma(\lambda_i) - \delta]. \quad (3.29)$$

For the second term in (3.28), according to the law of large numbers, we have

$$\frac{1}{n} \sum_{t=1}^n \mathbf{1}(Y_{a_i}^{(n)} \leq G_t < \gamma(\lambda_i) - \delta) = \mathbb{P}[Y_{a_i}^{(n)} \leq G < \gamma(\lambda_i) - \delta]. \quad (3.30)$$

The probability in (3.30) is upper bounded as

$$\mathbb{P}[Y_{a_i}^{(n)} \leq G < \gamma(\lambda_i) - \delta] \leq \mathbb{P}[Y_{a_i}^{(n)} < \gamma(\lambda_i) - \delta], \quad (3.31)$$

and according to Lemma 3.1,

$$\lim_{n \rightarrow \infty} \mathbb{P}[Y_{a_i}^{(n)} < \gamma(\lambda_i) - \delta] = 0. \quad (3.32)$$

Thus, (3.28) yields

$$\lim_{n \rightarrow \infty} \frac{1}{n} \sum_{t \in \mathbb{A}} I_i(G_t) \leq \mathbb{E}[I_i(G) | G \geq \gamma(\lambda_i) - \delta] \mathbb{P}[G \geq \gamma(\lambda_i) - \delta]. \quad (3.33)$$

Using (3.26) and considering the energy conservation,

$$n \sum_{i=1}^M \left[ \lambda_i + h \left( \frac{1}{n} \sum_{t \in \mathbb{A}} I_i(G_t) \right) \right] \leq n \sum_{i=1}^M \mathcal{E}_i \leq T\bar{W}. \quad (3.34)$$

This implies

$$\rho_n(T) = \frac{nRM}{T} \leq \frac{RM\bar{W}}{\sum_{i=1}^M \left[ \lambda_i + h \left( \frac{1}{n} \sum_{t \in \mathbb{A}} I_i(G_t) \right) \right]}, \quad (3.35)$$

and taking the limit while using (3.33)

$$\begin{aligned} \rho &= \lim_{n \rightarrow \infty} \lim_{T \rightarrow \infty} \rho_n(T) \leq \frac{MR\bar{W}}{\sum_{i=1}^M [\lambda_i + h(\mathbb{E}[I_i(G)|G \geq \gamma(\lambda_i)] \mathbb{P}[G \geq \gamma(\lambda_i)])]} \\ &= \frac{MR\bar{W}}{\sum_{i=1}^M [\mathbb{P}[G > \gamma(\lambda_i)] + \mathcal{E}_D(\gamma(\lambda_i))]} \end{aligned} \quad (3.36)$$

where in the last equality, we used (3.6) and (3.8), implying

$$h(\mathbb{E}[I_i(G)|G \geq \gamma(\lambda_i)] \mathbb{P}[G \geq \gamma(\lambda_i)]) = h(\hat{C}(\gamma(\lambda_i))) = \mathcal{E}_D(\gamma(\lambda_i)).$$

The upper bound is maximized when the total energy in the dominator is minimized.

Due to the convexity of the energy function,

$$\lambda^* = \arg \min_{\lambda} \mathbb{P}[G > \gamma(\lambda)] + \mathcal{E}_D(\gamma(\lambda)), \quad (3.37)$$

or  $\gamma^*$  as defined in (3.14) exist. So, the lower bound on the energy expenditure in all the packets 1 to  $M$  would be  $\mathcal{E}^*$  as defined in (3.15). Thus,

$$\rho \leq \frac{R\bar{W}}{\mathcal{E}^*}. \quad (3.38)$$

**Theorem 3.1** *For a fading channel with time-varying channel gain  $G_t$  and energy harvesting at the receiver, equipped with an unbounded battery, the optimum rate is*

given by

$$\rho = \frac{R\bar{W}}{\mathcal{E}^*},$$

where

$$\mathcal{E}^* = \Pr[G \geq \gamma^*] + \mathcal{E}_D(\gamma^*),$$

and

$$\gamma^* = \arg \min_{\gamma} \Pr[G \geq \gamma] + \mathcal{E}_D(\gamma),$$

and

$$\lambda^* = \Pr[G \geq \gamma^*].$$

*This rate is achieved using a variable-timing transmission protocol along with the channel selective sampling strategy with the threshold  $\gamma^*$ .*

### 3.5 Conclusion

We have considered a fading channel with energy harvesting at the receiver and an idealized situation at which the fading gain is known at the receiver before sampling the corresponding symbol. We have modeled the problem and obtained a tradeoff between the decoding energy and the characteristics of the channel. Based on this model, we proposed a channel selective sampling strategy and proved its optimality. We also characterized the optimal rate. The idealized model, used in this chapter, provides an upperbound to the systems in which the receiver samples must be used to estimate the channel.

Although the results in this chapter are reached using a deterministic model, we can extend them to the case with stochastic energy arrival when the codewords are long enough to experience the ergodic variation of the harvesting.

## Chapter 4

### Opportunistic Reception in a Multiuser Slow-Fading Channel with an Energy Harvesting Receiver

In this chapter, we focus on slow fading (block fading) channels where at the receiver side, the processing energy is harvested with limited rate. The transmitter sends data packets sequentially, where each packet carries coded data, separately encoded with a fixed rate. For each packet, the receiver should decide to sample and decode the packet or drop it and save the energy for next coming packets.

We derive an optimum policy to maximize the average throughput of the system. We show that a policy including a waiting period to charge the battery and a threshold for the channel gain to decide whether a packet should be decoded achieves the optimum rate with probability one. We then extend this result to multiuser systems with limited processing power at the receiver side. Serving the users with best channel gains increases the gap between the instantaneous capacity and the code rate, which in turn reduces the decoding energy. This saving in the processing energy requirement of each user allows the receiver to utilize its limited harvested energy to serve more users. In multiple access fading channels, serving the users with the strongest channels is known to provide a multiuser diversity gain [48–50] that can reduce the required *transmit power* or equivalently improve the overall rate. The results of this chapter extend the concept of multiuser diversity to reduce *processing power* at the receiver.

#### 4.1 System Model

We consider a communication system where  $K$  transmitters wish to communicate with one receiver with a slow-fading (block-fading) channel between transmitters and the receiver. Time is slotted, with each slot consisting of  $n$  channel uses. We will often



call a transmitted codeword a packet. We assume fixed-timing transmission strategy in which the packets are transmitted in every slot without any delay. Let the  $n$ -dimensional vectors  $\vec{X}_t^i$ ,  $\vec{Y}_t^i$ ,  $\vec{Z}_t^i$  respectively denote transmitted signal with average power  $p$ , received signal, and additive white Gaussian noise (AWGN) in slot  $t$  of the user  $i$ . The communication in slot  $t$  is modeled as

$$\vec{Y}_t^i = H_t^i \vec{X}_t^i + \vec{Z}_t^i, \quad (4.1)$$

where  $H_t^i$  denotes the complex channel coefficient of user  $i$  in time slot  $t$ . The channel coefficients  $H_t^i$  are fixed over the slot  $t$ , but vary independently both from one slot to another and from one user to another [47]. In addition,  $H_t^i$  for each user is known at the receiver causally, but not at the transmitters. We assume that receiver channel estimation takes negligible time compared to the packet length  $n$ . As there are  $K$  users, this requires  $n \gg K$ . The channel gain,  $G_t^i = |H_t^i|^2$ , is an i.i.d. sequence identical to  $G$ . The superscript  $i$  is dropped for single user system. During time-slot  $t$ , the capacity of the channel (4.1) for user  $i$  with channel gain  $G_t^i = g$ , is

$$C(g) = \log_2 \left( 1 + \frac{gp}{\sigma^2} \right). \quad (4.2)$$

We assume that the quantization bits are large enough so that the effect of the quantization error on the capacity is negligible.

When the sampling rate is  $\lambda$ , the energy per symbol (normalized energy) consumed by the receiver to reliably decode one packet will be

$$\mathcal{E} = \mathcal{E}(\lambda, g) = \nu\lambda + \mathcal{E}_D(\lambda, g). \quad (4.3)$$

We conclude for a given  $g$  that there is an optimal sampling rate  $\lambda^*(g)$  such that

$$\mathcal{E}^*(g) = \min_{R/C(g) < \lambda \leq 1} \mathcal{E}(\lambda, g) = \nu\lambda^*(g) + \mathcal{E}_D(\lambda^*, g) \quad (4.4)$$

is the minimum energy *per symbol period* required to decode a single rate  $R$  codeword.

Note that  $\mathcal{E}_D$  and as a result  $\mathcal{E}$  are increasing functions of packet length,  $n$ . We also observe that  $\mathcal{E}_D$  is a function of the code rate  $R$ . As a result, the optimum sampling rate and the minimum energy both depend on the code rate. We further assume that sampling and decoding are done jointly; any codeword which is sampled is decoded immediately and there is no option of keeping a codeword in the buffer and sampling the next codeword.

In this chapter, we use the above model to design opportunistic schemes for a single user and multiuser slow-fading channel to maximize the overall throughput or the expected number of users being served in every slot. We will show that choosing the users opportunistically will result in a larger capacity gap implying lower energy consumption at the receiver and ultimately better performance and multiuser diversity gain.

Similar to Chapter 3, here we assume that the energy  $\bar{W}$  arrives deterministically in every symbol period. We believe this is an appropriate model when codewords are transmitted in milliseconds and the coherence time of the energy harvesting process is on the order of minutes or hours. On the other hand, when codewords are much longer than the harvesting coherence time while the harvesting process is i.i.d., it was shown in Section 2.3 that system performance depends chiefly on the average harvesting rate.

## 4.2 Single User System

In this section, we assume a single transmitter communicating with a receiver through a slow-fading channel. According to the channel quality, the receiver may choose to sample and decode the transmitted codeword or it may save its stored energy for subsequent time slots in which the channel quality is better.

We assume that if a codeword is received at time  $t$ , an optimum number of samples,  $n\lambda^*(G_t)$  is taken and therefore, the optimum energy of  $\mathcal{E}^*(G_t)$  is used for sampling and decoding. The policy or scheme  $S$  is defined by the binary indicator  $I_S(t)$  such that  $I_S(t) = 1$  iff the codeword in the time slot  $t$  is received (sampled and decoded.). We note that  $S$  may or may not be designed such that the arrival energy  $n\bar{W}$  is always

sufficient for the energy requirements. For a given policy  $S$ , the stored energy  $U_t$  in the battery at time  $t$  satisfies

$$U_t = [U_{t-1} + n\bar{W} - I_S(t)n\mathcal{E}^*(G_t)]^+. \quad (4.5)$$

To mark whether the receiver has sufficient energy at time  $t$  to decode packet  $t$ , we need to evaluate the indicator  $\mathbf{1}(U_{t-1} + n\bar{W} \geq \mathcal{E}^*(G_t))$ . If this value is 1, the scheme is feasible at time  $t$ . For any feasible scheme  $S$ , the average throughput  $\rho_S$  is defined as

$$\rho_S = \liminf_{T \rightarrow \infty} \frac{R}{T} \sum_{t=1}^T I_S(t). \quad (4.6)$$

The objective is to maximize the average throughput (4.6) while ensuring the feasibility of the scheme at every time slot. We note that maximization of the rate  $\rho_S$  may be complex in that the optimal rate will be achieved by a packet processing policy that is a function of the receiver energy state.

#### 4.2.1 Throughput Optimization

We introduce the delayed-start threshold scheme,  $\tilde{S}$ , such that after some delay,  $T_1 = T_1(T) \in o(T)$ , packet  $t$  is decoded if the channel gain is above a threshold  $\tilde{\gamma}$  and the receiver has sufficient energy to sample and decode it. During the delay interval no packet is received and just the arrival energy is collected. We will see that choosing  $T_1 = T^{2/3}$  ensures that the throughput is not compromised. As the policy  $\tilde{S}$  will be designed to be feasible, the stored energy  $\tilde{U}$  under  $\tilde{S}$  can be specified recursively for  $t \geq 1$  by

$$\tilde{U}_t = \tilde{U}_{t-1} + n\bar{W} - n\mathcal{E}^*(G_t)I_{\tilde{S}}(t), \quad (4.7a)$$

$$I_{\tilde{S}}(t) = \begin{cases} 0 & t \leq T_1, \\ \mathbf{1}(G_t \geq \tilde{\gamma}, \tilde{U}_{t-1} + n\bar{W} \geq n\mathcal{E}^*(G_t)) & t > T_1. \end{cases} \quad (4.7b)$$

Furthermore, the threshold  $\tilde{\gamma}$  is selected such that

$$\mathbb{P}[G_t \geq \tilde{\gamma}] \mathbb{E}[n\mathcal{E}^*(G_t)|G_t \geq \tilde{\gamma}] = n\bar{W}. \quad (4.8)$$

Note that (4.8) has a unique solution as the left side is a non-increasing continuous function of  $\gamma$  and at some point it becomes equal to  $n\bar{W}$ . Note also that  $\tilde{\gamma}$  is such that  $C(\tilde{\gamma}) > R$ .

**Theorem 4.1** *For the policy  $\tilde{S}$  as defined in (4.7),*

$$\rho_{\tilde{S}} \geq R \mathbb{P}[G \geq \tilde{\gamma}] \quad w.p.1.$$

**Proof:** The throughput of the policy  $\tilde{S}$  is

$$\begin{aligned} \rho_{\tilde{S}} &= R \liminf_{T \rightarrow \infty} \frac{\sum_{t=T_1+1}^T I_{\tilde{S}}(t)}{T} \\ &= R \liminf_{T \rightarrow \infty} \left( \frac{T - T_1}{T} \right) \left( \frac{\sum_{t=T_1+1}^T I_{\tilde{S}}(t)}{T - T_1} \right). \end{aligned} \quad (4.9)$$

Since  $\lim_{T \rightarrow \infty} T_1/T = 0$ ,

$$\rho_{\tilde{S}} = R \liminf_{T \rightarrow \infty} \frac{\sum_{t=T_1+1}^T I_{\tilde{S}}(t)}{T - T_1}. \quad (4.10)$$

Defining  $V_t = n\bar{W} - n\mathcal{E}^*(G_t)I_{\tilde{S}}(t)$  as the energy increment in time  $t$ , the stored energy of the delayed start policy is

$$\tilde{U}_t = n\bar{W}T_1 + \sum_{\tau=T_1+1}^t V_{\tau}, \quad t > T_1. \quad (4.11)$$

To analyze the delayed start policy (4.7), we define

$$\hat{V}_t = n\bar{W} - n\mathcal{E}^*(G_t)\mathbf{1}(G_t > \tilde{\gamma}), \quad (4.12a)$$

$$\hat{U}_t = n\bar{W}T_1 + \sum_{\tau=T_1+1}^t \hat{V}(\tau), \quad t > T_1. \quad (4.12b)$$

We note that (4.7b) implies  $I_{\tilde{S}}(t) \leq \mathbf{1}(G_t \geq \tilde{\gamma})$ . Thus,  $V_t \geq \hat{V}_t$  and it follows from (4.11) and (4.12b) that  $\tilde{U}_t \geq \hat{U}_t$ . Thus (4.7b) implies

$$I_{\tilde{S}}(t) \geq \mathbf{1}(G_t \geq \tilde{\gamma}) \mathbf{1}\left(\hat{U}_{t-1} + n\bar{W} \geq n\mathcal{E}^*(G_t)\right). \quad (4.13)$$

Moreover, since  $\mathcal{E}^*(\gamma)$  is decreasing in  $\gamma$ , it follows that

$$\rho_{\tilde{S}} \geq \liminf_{T \rightarrow \infty} \frac{R}{T - T_1} \sum_{t=T_1+1}^T \mathbf{1}(G_t \geq \tilde{\gamma}) \hat{X}_t. \quad (4.14)$$

where  $\hat{X}_t$  denotes the indicator sequence

$$\hat{X}_t = \mathbf{1}\left(\hat{U}_{t-1} + n\bar{W} \geq n\mathcal{E}^*(\tilde{\gamma})\right). \quad (4.15)$$

We now show that the sequence  $\hat{X}_{T_1+1}, \dots, \hat{X}_T$  goes to one with probability one. It is sufficient to show that  $\lim_{T \rightarrow \infty} \mathbb{P}\left[\bigcap_{t=T_1+1}^T \{\hat{X}_t = 1\}\right] = 1$ . By the union bound,

$$\mathbb{P}\left[\bigcap_{t=T_1+1}^T \{\hat{X}_t = 1\}\right] \geq 1 - \sum_{t=T_1+1}^T \mathbb{P}[\hat{X}_t = 0]. \quad (4.16)$$

Note that since  $G_t$  is i.i.d.,  $\hat{V}_t$  is i.i.d.. From (4.8) and (4.12a), we observe that  $\mathbb{E}[\hat{V}_t] = 0$ . From (4.12a), (4.15) and noting that  $\hat{V}_t \in [n\bar{W} - n\mathcal{E}^*(\tilde{\gamma}), n\bar{W}]$ , and choosing  $T_1 = T^{2/3}$  it follows from the Chernoff-Hoeffding inequality that

$$\mathbb{P}[\hat{X}_t = 0] = \mathbb{P}\left[\sum_{\tau=T_1+1}^{t-1} \hat{V}_\tau \leq n\mathcal{E}^*(\tilde{\gamma}) - n\bar{W} - n\bar{W}T_1\right] \leq \epsilon_T(t) \quad (4.17)$$

where

$$\epsilon_T(t) = \exp \frac{-2[\bar{W}(1 + T^{2/3}) - \mathcal{E}^*(\tilde{\gamma})]^2}{(t - T^{2/3} - 1)\mathcal{E}^*(\tilde{\gamma})^2}. \quad (4.18)$$

It then follows from (4.16) that

$$\begin{aligned} \lim_{T \rightarrow \infty} \mathbb{P} \left[ \bigcap_{t=T_1+1}^T \{\hat{X}_t = 1\} \right] &\geq \lim_{T \rightarrow \infty} 1 - \sum_{t=T_1+1}^T \mathbb{P} [\hat{X}_t = 0] \\ &\stackrel{a}{\geq} \lim_{T \rightarrow \infty} 1 - (T - T_1) \epsilon_T(T) = 1, \end{aligned} \quad (4.19)$$

where (a) is true since  $\epsilon_t(T) \leq \epsilon_T(T)$ . Note that  $\lim_{T \rightarrow \infty} \epsilon_T(T) = 0$ . As a result,  $\hat{X}_t \rightarrow 1$  w.p.1. In (4.14), the event  $G_t \geq \tilde{\gamma}$  and the indicator  $\hat{X}_t$  are independent. It is straightforward to see that (4.14) can be rewritten as

$$\rho_{\tilde{S}} \geq R \left( \liminf_{T \rightarrow \infty} \frac{1}{T - T_1} \sum_{t=T_1+1}^T \mathbf{1}(G_t \geq \tilde{\gamma}) \right) \left( \liminf_{T \rightarrow \infty} \hat{X}_t \right). \quad (4.20)$$

According to the strong law of large numbers, we have

$$\liminf_{T \rightarrow \infty} \frac{1}{T - T_1} \sum_{t=T_1+1}^T \mathbf{1}(G_t \geq \tilde{\gamma}) = \mathbb{P}[G \geq \tilde{\gamma}] \quad \text{w.p.1}, \quad (4.21)$$

and therefore,  $\rho_{\tilde{S}} \geq R \mathbb{P}[G \geq \tilde{\gamma}]$  w.p.1.  $\square$

It can be seen that the key is that after the delayed start, the available energy in the battery is lower bounded by the process in which packet  $t$  is decoded if just the threshold constraint  $G_t \geq \tilde{\gamma}$  is satisfied. This in turn is lower bounded by the stored energy of the process in which the energy to process a packet is upper bounded by  $\mathcal{E}^*(\tilde{\gamma})$ , the energy needed for a packet just meeting the threshold channel quality constraint. We show that the collected energy by time  $T_1$  is sufficient to decode almost all packets after  $T_1$  with the channel above the threshold.

We now show that no scheme can do better than this opportunistic selection scheme. Assume an arbitrary feasible scheme  $S$  decodes  $M_T$  packets transmitted in  $T$  time slots. It is easy to see that a noncausal scheme that decodes the  $M_T$  packets with the best  $M_T$  channel gains would require less total energy. We use this observation to prove the following claim.

**Theorem 4.2** *If  $S$  is a feasible receiver that decodes  $M_T$  packets in  $T$  slots such that*

$\lim_{T \rightarrow \infty} M_T/T = \beta$ , then  $\beta \leq \mathbb{P}[G_t \geq \tilde{\gamma}]$ .

**Proof:** To prove this claim, we will need to define  $b_T = \lfloor \beta T \rfloor$  and  $\gamma(\beta)$  such that it satisfies

$$\bar{F}_{G_t}(\gamma(\beta)) = \mathbb{P}[G_t \geq \gamma(\beta)] = \beta. \quad (4.22)$$

In addition, we use  $Y_1^{(T)}, Y_2^{(T)}, \dots, Y_T^{(T)}$  to denote the random variables  $G_1, G_2, \dots, G_T$  sorted in decreasing order and we define  $Y_{\beta, T} \triangleq Y_{b_T}^{(T)}$ . According to Lemma 3.1,  $Y_{\beta, T}$  converges in distribution to  $\gamma(\beta)$  as  $T \rightarrow \infty$ .

Over  $T$  slots, the scheme  $S$  has total receiver energy consumption satisfying  $\mathcal{E}^{\text{total}} \leq T\bar{W}$ . Let  $B_T \subseteq \{1, 2, \dots, T\}$  denote the subset of time slots  $t$  with the best  $b_T - 1$  channels. The normalized energy required to decode this subset of packets is

$$\mathcal{E}_T^* = \sum_{t \in B_T} \mathcal{E}^*(G_t) \leq \mathcal{E}^{\text{total}} \leq T\bar{W}. \quad (4.23)$$

We define the indicator  $I_{\beta, T}(t)$  to equal 1 if  $t \in B_T$  and zero otherwise. This permits us to write

$$\bar{W} \geq \frac{1}{T} \sum_{t=1}^T \mathcal{E}^*(G_t) I_{\beta, T}(t). \quad (4.24)$$

Since the  $G_t$  are i.i.d., the random variables  $\mathcal{E}^*(G_t) I_{\beta, T}(t)$  are identically distributed. Thus, taking the expectation of (4.24), we obtain

$$\bar{W} \geq \frac{1}{T} \sum_{t=1}^T \mathbb{E}[\mathcal{E}^*(G_t) I_{\beta, T}(t)] = \mathbb{E}[\mathcal{E}^*(G_t) I_{\beta, T}(t)]. \quad (4.25)$$

Using  $F_{\beta, T}(y)$  denote the CDF of  $Y_{\beta, T}$ , we can write

$$\bar{W} \geq \int \mathbb{E}[\mathcal{E}^*(G_t) I_{\beta, T}(t) | Y_{\beta, T} = y] dF_{\beta, T}(y). \quad (4.26)$$

where

$$\mathbb{E}[\mathcal{E}^*(G_t)I_{\beta,T}(t)|Y_{\beta,T}=y] = \mathbb{E}[\mathcal{E}^*(G_t)|Y_{\beta,T}=y, I_{\beta,T}(t)=1]\mathbb{P}[I_{\beta,T}(t)=1|Y_{\beta,T}=y]. \quad (4.27)$$

We observe that symmetry implies that

$$\mathbb{P}[I_{\beta,T}(t)=1|Y_{\beta,T}=y] = \frac{b_T-1}{T} \quad (4.28)$$

since each  $G_t$  is equally likely to be one of the  $b_T - 1$  channel gains better than  $Y_{\beta,T}$ . In addition,

$$\mathbb{E}[\mathcal{E}^*(G_t)|Y_{\beta,T}=y, I_{\beta,T}(t)=1] = \mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq y]. \quad (4.29)$$

Thus (4.26), (4.27), (4.28) and (4.29) imply

$$\bar{W} \geq \int \mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq y] \left( \frac{b_T-1}{T} \right) dF_{\beta,T}(y). \quad (4.30)$$

As  $T \rightarrow \infty$ ,  $b_T/T \rightarrow \beta$  and, by Lemma 3.1,  $F_{\beta,T}(y)$  converges to a unit step at  $\gamma(\beta)$ . Letting  $T \rightarrow \infty$  in (4.30), we obtain

$$\bar{W} \geq \beta \mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq \gamma(\beta)]. \quad (4.31)$$

It follows from (4.22) that

$$\bar{W} \geq \mathbb{P}[G_t \geq \gamma(\beta)] \mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq \gamma(\beta)]. \quad (4.32)$$

Now we recall that  $\mathbb{P}[G_t \geq y] \mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq y]$  is a non-increasing function of  $y$ . Thus, (4.8) and (4.32) imply  $\tilde{\gamma} \leq \gamma(\beta)$ . Moreover, since  $\mathbb{P}[G_t \geq y]$  is decreasing in  $y$ ,

$$\beta = \mathbb{P}[G_t \geq \gamma(\beta)] \leq \mathbb{P}[G_t \geq \tilde{\gamma}]. \quad (4.33)$$

□



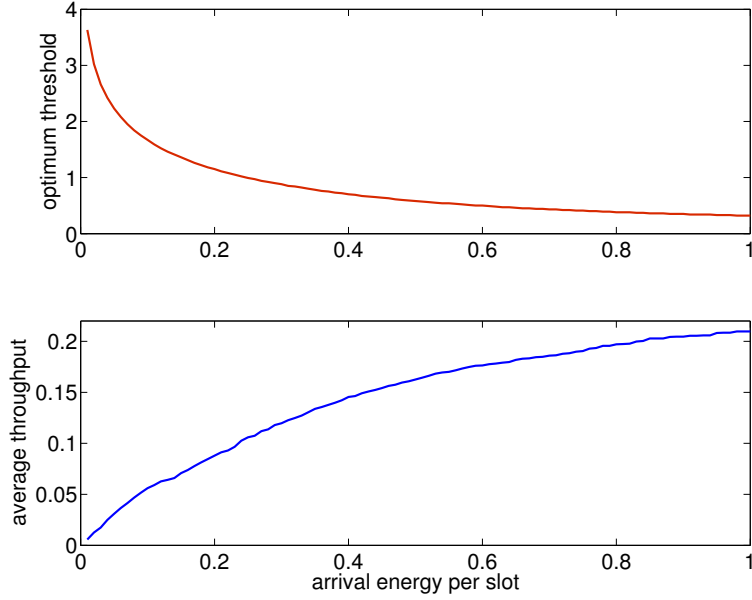


Figure 4.1: Optimum threshold  $\tilde{\gamma}$  and optimum throughput versus  $n\bar{W}$  for  $R = 0.3$  bits/s/Hz,  $\nu = 1.69$ ,  $\eta = 0.42$ .

We note that Theorem 4.2 implies that  $\rho_S \leq R \mathbb{P}[G \geq \tilde{\gamma}]$  for any scheme  $S$ .

#### 4.2.2 Single User Performance/ Self-imposed Threshold

As depicted in Fig. 4.1, the threshold  $\tilde{\gamma}$  decreases with increasing  $\bar{W}$  and as a result the throughput improves. For this figure, we used the decoding energy model in [27] as

$$\mathcal{E}_D = \left( \frac{\eta}{1 - R/\lambda C(g)} \right) \log \left( \frac{1}{1 - R/\lambda C(g)} \right)$$

where  $\eta$  is a coefficient and  $C(g)$  is defined in (4.2). For fixed code rate  $R$ , we define the *service rate* as the average number of packets decoded per slot. In Fig. 4.2, we plot the service rate as a function of threshold  $\tilde{\gamma}$ . When the receiver employs the threshold  $\gamma > \tilde{\gamma}$ , the arrival energy per slot exceeds the average energy required per slot and the resulting service rate is  $\mathbb{P}[G > \gamma]$ , the probability that the channel is above threshold. On the other hand, when  $\gamma < \tilde{\gamma}$ , the energy demand is larger than the arrival energy in a slot; even if the channel gain is above the threshold, the packet may not be decoded due to insufficient energy at the receiver. In this case, one might guess the service rate

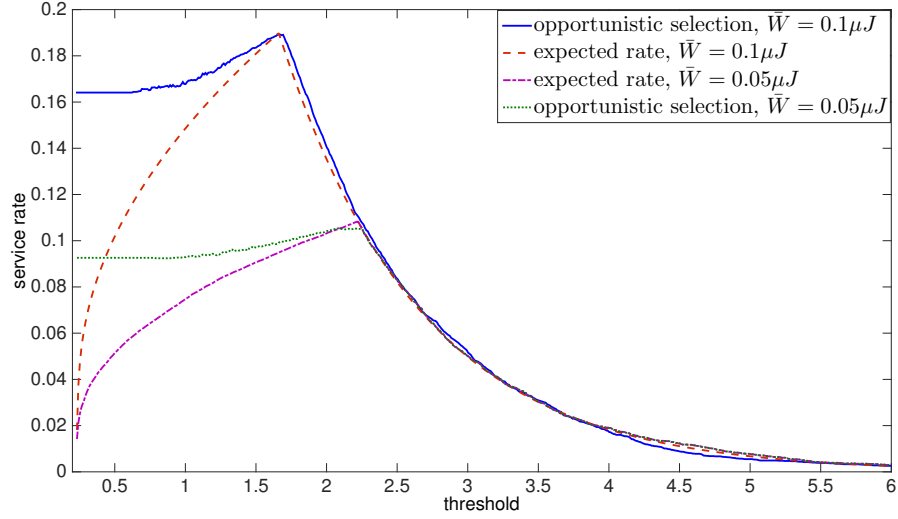


Figure 4.2: Comparing the achievable service rate of the opportunistic selection with what suggested in (4.34),  $n = 10^6$ .

would be  $\bar{W}/\mathbb{E}[\mathcal{E}^*(G)|G \geq \gamma]$ . In general, it appears that the service rate of

$$\min \left\{ \mathbb{P}[G \geq \gamma], \frac{\bar{W}}{\mathbb{E}[\mathcal{E}^*(G)|G \geq \gamma]} \right\} \quad (4.34)$$

would be achievable. However, when  $\gamma < \tilde{\gamma}$ , the available energy in the battery is typically close to zero. In this case, whenever the stored energy reaches a level sufficient to sample and decode a packet, the packet is sampled immediately. So the receiving decision in slot  $t$  is strongly dependent on the energy level of the battery in that slot. And since the energy level is often low, packets are decoded very selectively. That is, we call this phenomenon a *self-imposed threshold*. That is, when the energy level is low, only packets transmitted through very good channels are decoded. Due to this fact, we see in Fig. 4.2 that the achievable rate is much higher than the rate suggested by (4.34). Evaluation of the service rate in this situation requires an analysis of the continuous-value Markov chain  $U_t$  with state dependent reward  $I_S(t)$ . The complexity of this analysis highlights the value of the outer bound result of Theorem 4.2.

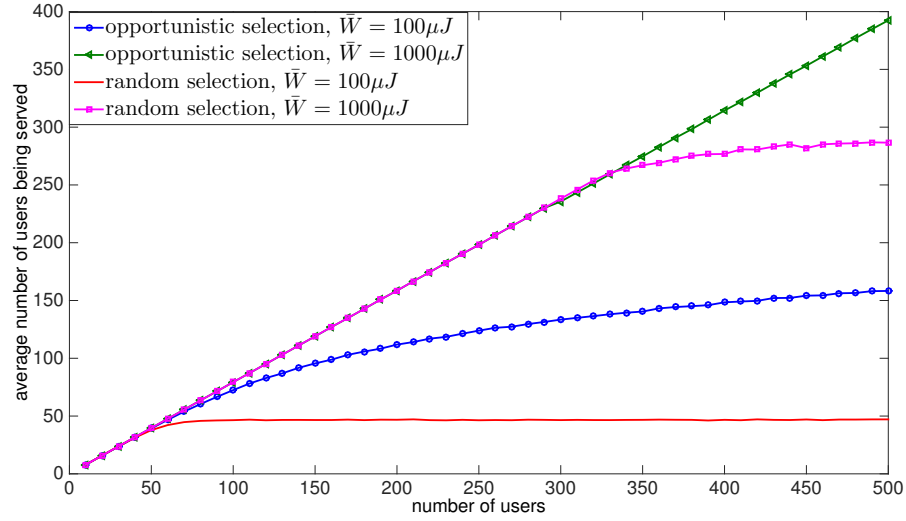


Figure 4.3: Comparison of the number of users being served in the opportunistic and random selection for two  $\bar{W} = 100 \mu J$  and  $\bar{W} = 1000 \mu J$  when  $R = 0.3$  bits/s/Hz,  $\nu = 1.69$ ,  $\eta = 0.42$ ,  $n = 10^6$ .

### 4.3 Multiuser System

Assume there are  $K$  users transmitting at a fixed rate  $R$  and fixed average power  $p$  over orthogonal Rayleigh fading channels where the channel gain of user  $i$  in slot  $t$  is denoted as  $G_t^i$ . Energy harvesting limits the energy available at the receiver so that it may not be possible to decode all the packets. The objective is to maximize the sum-rate as defined as

$$\rho_{S_K} = \liminf_{T \rightarrow \infty} \frac{R}{T} \sum_{t=1}^T \sum_{i=1}^K I_S^{(i)}(t), \quad (4.35)$$

where  $I_S^{(i)}(t) = 1$  if the user  $i$ 's packet transmitted at time  $t$  is decoded.

We observe that the  $K$ -user system can be viewed as a faster-paced single user system. In particular, imagine that a slot is divided into  $K$  minislots such that the receiver harvests energy  $n\bar{W}/K$  in each minislot and user  $k$  transmits in minislot  $k$  of each slot. With the minislots as the unit of time, we have a system with one packet transmitted each time unit and energy  $n\bar{W}/K$  harvested in each time unit. From our single-user analysis, we know that we cannot do better than a threshold policy such that

we attempt to decode every packet with channel gain above a threshold  $\tilde{\gamma}_K$  satisfying

$$\mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq \tilde{\gamma}_K]\mathbb{P}[G \geq \tilde{\gamma}_K] = \bar{W}/K, \quad (4.36)$$

which results in the rate

$$\rho_{S_K} = R K \mathbb{P}[G \geq \tilde{\gamma}_K] = \frac{R \bar{W}}{\mathbb{E}[\mathcal{E}^*(G_t)|G_t \geq \tilde{\gamma}_K]}. \quad (4.37)$$

For a fixed code rate,  $R$ , and a fixed sampling rate,  $\lambda$ , the capacity gap corresponding to the selected users, increases with  $K$ , which reduces the required decoding energy. The sampling rate is selected such that the minimum energy is consumed to sample and decode each packet. This implies serving more users under a limited energy harvesting rate leading to the maximum sum-rate. Fig. 4.3 compares the number of users being served in this opportunistic scheme with the one for random selection for two choices of  $\bar{W}$  as  $K$  is growing. It can be seen that while the service rate of the random selection is fixed with  $K$ , it is growing for the opportunistic selection. We note that while the threshold grows with the number of users, the number of the selected channels above the threshold also grows with  $K$ . Therefore the number of users being served grows with  $K$ .

Note that the sum-rate also depends on the code rate,  $R$ . Increasing  $R$  reduces the capacity gap, thus increasing the processing energy while also increasing the number of information bits communicated when a packet is decoded. As it can be seen in Fig. 4.4 for larger number of users, increasing the processing energy appears to be the more dominant effect and the system throughput suffers as the code rate is increased.

#### 4.4 Conclusions

In this chapter, we exploited the variation of channel over both time and users to reduce the processing power per decoded packet at the receiver. This is in contrast with conventional systems in which this channel variation is used to save transmit power or increase the packet code rate. This saving in the receiver processing power is

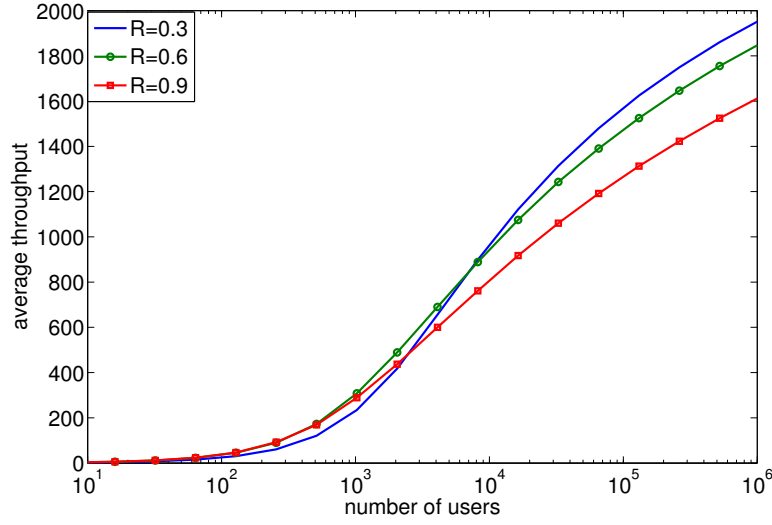


Figure 4.4: The sum-rate versus the number of users for  $R = 0.3, 0.6, 0.9$  bits/s/Hz when  $\eta = 0.1$ ,  $\nu = 2.5 \times 10^{-4}$  and  $\bar{W} = 100 \mu J$ ,  $n = 10^8$ .

particularly useful in energy harvesting systems where the rate of energy arrival at the receiver may constrain the packet decoding rate.

We have shown that an energy-constrained receiver can benefit from a multiuser diversity effect in which the users with unusually good channels enable increased efficiency in the receiver sampling and decoding. In our numerical examples, we observe that this effect is not especially sensitive to the choice of transmitter code rate. This insensitivity arises primarily because users operating at reduced code rates enable the energy-constrained receiver to process more users.

## Chapter 5

### Hybrid ARQ in Block-Fading Channels with an Energy Harvesting Receiver

In this chapter, we propose using ARQ to reduce the processing energy at the receiver, and therefore increase the overall throughput, when the receiver harvests energy with a limited rate. Conventionally, ARQ is used to improve the throughput and reliability, as well as to reduce the delay [51–54] in a fading channel with limited *transmit power*. In this chapter, however, our focus is on the receiver, where *processing power* is limited. ARQ can help to improve the *capacity gap*, which is defined as the gap between the code rate and the mutual information between transmitter and receiver. We will focus on the most popular ARQ schemes: (1) IR-HARQ, which sends additional parity symbols in each retransmission, (2) Repetition-HARQ, which repeats sending the same coded packet in each retransmission. We observe that depending on the parameters of the problem, Repetition-HARQ can perform better than IR-HARQ. This is in contrast to systems without any constraint on receiver processing power. This is due to the fact that the decoding energy is a decreasing function of the capacity gap and an increasing function of the code-length. IR-HARQ yields a better capacity gap, but increases the code-length, while Repetition-HARQ offers less improvement in the capacity gap, but does not increase the effective code-length. We also propose the scheme of IR-HARQ with subset selection which optimizes the subset of packets contributing to the decoding.

#### 5.1 System Model

We consider a point-to-point communication system with a block-fading channel between the transmitter and the receiver. Time is slotted, where each slot  $t$  consists of

$n$  channel uses. Let the  $n$ -dimensional vectors  $\vec{X}_t$ ,  $\vec{Y}_t$ ,  $\vec{Z}_t$  respectively denote the transmitted signal, the received signal, and the additive white Gaussian noise (AWGN) in slot  $t$ . We consider a narrow-band block fading environment, where the channel coefficient  $H_t \in \mathbb{C}$  is fixed during the time slot  $t$ , but varies independently from one slot to another. The communication in slot  $t$  is modeled as

$$\vec{Y}_t = H_t \vec{X}_t + \vec{Z}_t. \quad (5.1)$$

$H_t$  is known at the receiver causally, but it is not available at the transmitter. The average transmit power is  $p$  and the AWGN noise is i.i.d. across time with average zero and variance  $\sigma^2$ . The channel gain,  $G_t = |H_t|^2$ , is an i.i.d. sequence, with distribution identical to random variable  $G$ .

During the time-slot  $t$  with channel gain  $G_t$ , the maximum mutual information between the transmit signal and the received signal at each channel use (symbol period), denoted by  $C_t$ , is equal to  $C_t = C(G_t)$ , where

$$C(g) = \log_2 \left( 1 + \frac{gp}{\sigma^2} \right). \quad (5.2)$$

For a set of slots  $\mathcal{T} = \{1, \dots, k\}$ , we can show that the maximum mutual information between transmitted signals  $\{\vec{X}_t\}_{t \in \mathcal{T}}$  and received signals  $\{\vec{Y}_t\}_{t \in \mathcal{T}}$  is equal to  $kn \sum_{t \in \mathcal{T}} C_t$ .

In general, the processing energy at the receiver includes the energy consumption of sampler, decoder and RF front end. In this chapter, we only focus on the decoding energy, assuming the energy consumption by the other components is fixed. We follow the decoding energy model presented in the Chapter 1. That is, for a super-block containing  $L$  blocks of data of length  $n$ , the decoding energy  $E_D$  is

$$E_D = Ln\mathcal{E}_D(\delta) = Ln\mathcal{E}_D(1 - R/C), \quad (5.3)$$

In this work, we assume that the energy arrival process is deterministic, with a constant rate of  $n\bar{W}$  Joules per time slot, which is available at the beginning of each

slot. The receiver is equipped with an unlimited-capacity battery to store the energy.

## 5.2 ARQ Schemes

In this chapter, we consider various forms of ARQ schemes. Here we briefly describe these alternatives.

### Classic ARQ

In this scheme, the transmitter uses a Gaussian codebook of length  $n$  and rate  $R$  to send  $nR$  bits in slot  $t$ , using  $n$ -symbol *coded packets* in  $\vec{X}_t$ . Each coded packet requires one slot time for transmission. The receiver may decide to decode the message from  $\vec{Y}_t$ . In this case, the receiver sends back an ACK to let the transmitter know its decision. Otherwise, if the receiver decides not to decode the message, it drops  $\vec{Y}_t$  and sends a NAK to the transmitter. We assume the ACK/NAK is received by the receiver reliably. If the transmitter receives an ACK, it starts sending a new message. Otherwise, it retransmits the previous signal again in slot  $t + 1$ , i.e.,  $\vec{X}_{t+1} = \vec{X}_t$ . The receiver receives  $\vec{Y}_{t+1}$  through a statistically independent new channel, and decides to whether decode the message from  $\vec{Y}_{t+1}$ , or drops it, and asks for a retransmission by returning a NAK. The retransmissions continue until the message is decoded from the last received packet. We consider no delay constraint for the system, so the NAK feedback and retransmission can be repeated until the message is decoded reliably. The receiver decision on whether to decode the message or to ask for a retransmission is based on the energy required for decoding and also the available accumulated energy in the battery. We note that there are two necessary conditions for decoding: (i) the rate  $R$  needs to be less than maximum mutual information in the last time slot where decoding occurs and (ii) the available energy in the battery must exceed the required decoding energy. We assume that  $n$  is large enough that decoding can be executed with arbitrary small probability of error. We also ignore the energy consumption of the ACK/NAK signal as well as any delay or unreliability of the feedback channel.



### Repetition-HARQ

In Repetition-HARQ, the same coded packet is repeatedly transmitted and the receiver combines the received copies through MRC until the message is decoded reliably.

### IR-HARQ

In contrast to the previous two methods, IR-HARQ avoids retransmitting the same codeword in every retransmission. In this scheme, the transmitter utilizes a Gaussian codebook, with  $2^{nR}$  very long codewords. To transmit a message of  $nR$  bits, the transmitter chooses the corresponding codeword, and sends the first  $n$  symbols of the codeword. If the receiver asks for retransmission by returning a NAK, the transmitter sends the second  $n$  symbols of that codeword and so on. Here, the receiver employs the contributions of all received packets to decode a message; therefore, the additional redundancy added at each step helps the receiver to decode the message successfully. This model is contrast to [55] where in each retransmission the transmitter sends out a punctured code with a different power level while the bits are selected probabilistically. We note that if we truncate all the codewords, and only keep the first  $kn$  symbols of each codeword, for some positive integer  $k$ , the new codebook is equivalent to a Gaussian codebook of rate  $R/k$ .

#### 5.2.1 Objective Function

Assume the number of decoded messages by slot  $t$  is denoted by  $m(t)$ . The throughput,  $\rho$ , is defined as the average number of successfully decoded information bits per symbol period:

$$\rho(R, \bar{W}) = \liminf_{t \rightarrow \infty} R \frac{m(t)}{t}. \quad (5.4)$$

Note that the throughput is a function of  $R$  and  $\bar{W}$ . In this work, the objective is to design the decoding and retransmission policy in order to maximize the throughput while the energy causality constraint is satisfied. Energy causality guarantees that the total energy used at the receiver up to slot  $t$  must be less than  $tn\bar{W}$  for all  $t$ .

### 5.3 Decoding Energy

The limited energy arrival rate forces the receiver to keep requesting retransmissions to not only increase the capacity gap and reduce the decoding energy, but also to collect and save enough energy for decoding in the battery. Therefore, the retransmission time, and thus the average throughput, is a function of the decoding energy. In order to maximize the throughput, it is therefore important to have a close look at the decoding energy of each of these ARQ schemes.

Assume that in an ARQ scheme, the transmission for  $nR$  information bits uses slots  $\mathcal{T} = \{1, 2, \dots, K\}$  (i.e.  $K - 1$  retransmissions), for some positive integer  $K$ . The receiver decodes the data at the end of slot  $K$ . We note that in general  $K$  is random, a function of channel realizations and energy of the battery. Here we discuss the decoding energy of each scheme for a given  $K = k$  and channel gains  $g_1, \dots, g_k$ .

For simplicity of exposition, we define a short-hand notation

$$f_R(C) \triangleq n\mathcal{E}_D(1 - \frac{R}{C}), \quad (5.5)$$

which is a decreasing function of  $C$ .

#### Classic ARQ

Since in this scheme, the receiver drops all earlier packets and decodes the data from the very last received packet, the decoding energy is equal to  $f_R(C_k)$ , where  $C_k = C(g_k)$  is the maximum mutual information per channel use during slot  $k$ .

#### Repetition-HARQ

The message is encoded into a length- $n$  codeword (coded packet) and retransmitted until the receiver decides to decode the message. From maximum ratio combining, we know that the sufficient statistics for decoding the message from  $\{\vec{Y}_t\}_{t=1}^k$  is  $\sum_{t=1}^k h_t^* \vec{Y}_t$ , where  $h_t$  denotes the channel coefficients in slots  $t = 1, \dots, k$ . Therefore, the maximum

mutual information between the transmitter and the receiver is

$$C\left(\sum_{t=1}^k g_t\right) = \log_2\left(1 + \frac{p}{\sigma^2} \sum_{t=1}^k g_t\right), \quad (5.6)$$

and the decoding energy is equal to

$$f_R\left(C\left(\sum_{t=1}^k g_t\right)\right). \quad (5.7)$$

### IR-HARQ

As we explained in subsection 5.1, the mutual information between transmitted signals  $\{\vec{X}_t\}_{t \in \mathcal{T}}$  and received signals  $\{\vec{Y}_t\}_{t \in \mathcal{T}}$  in  $k$  slots is equal to  $n \sum_{t=1}^k C_t = n \sum_{t=1}^k C(g_t)$ .

The required energy to decode the first message after  $k$  transmissions is

$$kn\mathcal{E}_D\left(1 - \frac{R/k}{\sum_{t=1}^k C_t/k}\right) = kf_R\left(\sum_{t=1}^k C_t\right). \quad (5.8)$$

The fact that the energy of decoding is not only a function of the capacity gap, but also a function of the code length, has some interesting implications. These implications lead to some results that are contrary to observations that are valid for ARQ systems with no constraint on the decoding energy.

- Let us compare the decoding energy of the IR-HARQ and Repetition-HARQ systems. Since  $C(\cdot)$  is concave,  $C\left(\sum_{t=1}^k g_t\right) \leq \sum_{t=1}^k C(g_t)$ . Since  $f_R(C)$  is a decreasing function of  $C$ , therefore  $f_R\left(C\left(\sum_{t=1}^k g_t\right)\right) \geq f_R\left(\sum_{t=1}^k C(g_t)\right)$ . However, it is not enough to conclude that the decoding energy of IR-HARQ is less than that of Repetition-HARQ. The reason is that in Repetition-HARQ, the effective length of the code after MRC is fixed and equal to one slot, i.e.  $n$ . However, in IR-HARQ, the length of a codeword is equal to the length of  $k$  slots, i.e.  $kn$ . Recall that the decoding energy linearly scales with  $k$ . This is reflected in the prefactor  $k$  in the decoding energy of IR-HARQ in (5.8). As a result, depending on the parameters of the problem, either the Repetition-HARQ or IR-HARQ would perform better. This is in contrast with the case where decoding energy is

not limited, and IR-HARQ always performs better than Repetition-HARQ [52].

- Let us now focus on the decoding energy for IR-HARQ itself. Every packet received at slot  $\hat{t} \in \{1, \dots, k\}$  increases  $\sum_t C_t = \sum_t C(g_t)$  by  $C(g_{\hat{t}})$ , improving the capacity gap. On the other hand, it also increases the code length by one slot. Therefore, it is not clear if the received packet during  $\hat{t}$  indeed reduces the decoding energy. For example, if  $g_{\hat{t}}$  is very small, the contribution of the packet sent in time slot  $\hat{t}$  in reducing the capacity gap is negligible, however it increases the length of the code by one slot. Therefore, it increases the decoding energy. In contrast, in HARQ without an energy constraint at the receiver, every single observation, as long as the corresponding channel gain is not equal to zero, is helpful, no matter how small it is. We suggest the scheme of *IR-HARQ with subset selection* which is explained next.

### IR-HARQ with subset selection

Because the decoding energy is a function of both the capacity gap and the code-length, it is sometimes better in IR-HARQ to ignore some received packets in decoding the message. Thus, we propose that the receiver chooses a subset  $\mathcal{T} \subset \{1, \dots, k\}$  of received packets to decode the message. If the receiver decodes the message from the received packets in slots  $\mathcal{T}$ , then the decoding energy is equal to

$$n|\mathcal{T}|\mathcal{E}_D \left( 1 - \frac{R/|\mathcal{T}|}{\sum_{t \in \mathcal{T}} C_t/|\mathcal{T}|} \right) = |\mathcal{T}|f_R \left( \sum_{t \in \mathcal{T}} C_t \right). \quad (5.9)$$

Therefore, the minimum energy of decoding is equal to

$$\min_{\mathcal{T} \subset \{1, \dots, k\}} |\mathcal{T}|f_R \left( \sum_{t \in \mathcal{T}} C_t \right). \quad (5.10)$$

Note that in IR-HARQ with subset selection, the transmission scheme is the same as IR-HARQ, the only difference is that in the process of decoding or calculating the decoding energy, we use only a subset of the transmitted packets.

**Remark:** We note that to find the optimum subset of slots in (5.10), we do not need

to go through all subsets of  $\{1, 2, \dots, k\}$ . Consider a permutation  $\boldsymbol{\pi} = (\pi(1), \dots, \pi(k))$  of the set  $\{1, 2, \dots, k\}$ , such that  $C_{\pi(1)} \geq C_{\pi(2)} \geq \dots \geq C_{\pi(k)}$ . Then, among all subsets  $\mathcal{T}$ , with cardinality  $\ell$ , for some  $\ell$ ,  $1 \leq \ell \leq k$ , the subset  $\{\pi(1), \dots, \pi(\ell)\}$  of the received packets achieves the largest mutual information and as a result requires the least decoding energy. Therefore, the minimum energy of decoding is equal to

$$\min_{\ell \in \{1, \dots, k\}} \ell f_R \left( \sum_{i=1}^{\ell} C_{\pi(i)} \right). \quad (5.11)$$

The complexity of the above optimization is linear in  $k$ .

#### 5.4 Decision Policy

To derive the the optimum decision policy, maximizing the throughput, one option is to use *dynamic programming* to develop a recursive optimization problem [56], which yields the best decision at each slot based on the collected packets, stored energy in the battery, and the expected performance in the future. However, the energy left in the battery after each decoding couples the decision process for different messages making the problem non-trivial.

As an alternative solution, here we propose a scheme in which the receiver decides to decode a packet if (1) it has enough energy accumulated in the battery to decode it, (2) the required energy to decode is less than a threshold  $\gamma$ . Otherwise a retransmission is requested. The threshold  $\gamma$  is selected such that the throughput is maximized. Note that the condition (2) presumes that the capacity gap is positive. Otherwise, the decoding energy is unboundedly large. In the next section, we present simulation results for this scheme.

In the threshold-based classic ARQ, the retransmission is continued until the receiver receives a packet through a good enough channel gain which requires decoding energy less than a threshold  $\gamma$ , and at the same time, the battery has enough energy to decode the message. Since in classic ARQ in each slot, the receiver drops the previous packets, this problem reduces to another problem, in which in each slot, the transmitter sends a new coded packet, and the receiver opportunistically decodes the packet. In Chapter

4 (and also [23],) we proved that a threshold-based algorithm, concatenated with a waiting time at the beginning achieves the optimum performance under this setting<sup>1</sup>. The optimum throughput  $\rho^*$  is equal to

$$\rho^* = R \Pr[f_R(C_t) \leq \gamma^*] \quad w.p.1, \quad (5.12)$$

where the optimum threshold  $\gamma^*$  is the solution to the following equation

$$\Pr[f_R(C_t) \leq \gamma^*] \mathbb{E}[f_R(C_t) | f_R(C_t) \leq \gamma^*] = n\bar{W}.$$

The threshold-based algorithm we propose here is motivated by this result.

## 5.5 Simulation Results

In this section, we compare various ARQ schemes applying threshold-based decision. Here, we let  $\mathcal{E}_D(\delta) = (1/\delta) \log(1/\delta)$  [27, 31] and  $p/\sigma^2 = 1$ . The channel coefficient  $H$  has a Rayleigh distribution with second moment 1.

Figures 5.1 and 5.2 compare throughput versus  $n\bar{W}$  for all four schemes for  $R = 3$  bit/s/Hz and  $R = 0.5$  bit/s/Hz, respectively. We have the following observations:

- As discussed, depending on the values of  $R$  and  $\bar{W}$ , the Repetition-HARQ can perform either better or worse than IR-HARQ with subset selection.
- For large values of  $\bar{W}$ , the system cares less about the energy, so it seeks to increase the mutual information to exceed  $R$  as early as possible to start decoding the message. Since IR-HARQ improves the overall mutual information faster than Repetition-HARQ, it performs better for large  $\bar{W}$ . For small  $\bar{W}$ , energy is the main constraint and thus Repetition-HARQ performs better.
- Since both the effective code length and the capacity gap in classic ARQ are less than or equal to those of IR-HARQ, it is possible that the classic ARQ outperforms

---

<sup>1</sup>In contrast to Chapter 4 where the threshold was on the channel gain, here we put the threshold on the decoding energy.

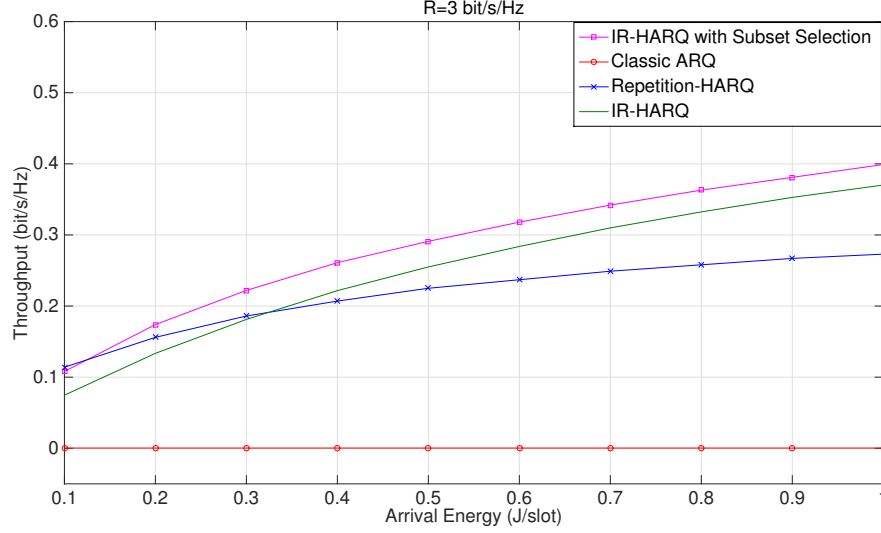


Figure 5.1: Throughput versus the arrival energy rate of  $n\bar{W}$  for the code rate of  $R = 3$  bits/s/Hz,  $n = 10^6$ .

IR-HARQ, but it never outperforms IR-HARQ with subset selection. Subset selection always results in an effective code length that minimizes the decoding energy.

- IR-HARQ performs poorly when  $R$  is small because the code length grows quickly with retransmissions.

In ARQ-based schemes, the transmitter relies on the channel statistics to set the rate  $R$ , such that the overall throughput is maximum. One important question is whether, Repetition-HARQ can still perform better than IR-HARQ for the optimum choice of  $R$ . Figures 5.3 and 5.4 compare throughput versus  $R$  for all four schemes for two different energy arrival rates of  $n\bar{W} = 0.3$  J/slot and  $n\bar{W} = 2$  J/slot, respectively. We can see in Fig. 5.3, that for  $n\bar{W} = 0.3$  J/slot, the maximum throughput of Repetition-HARQ is  $\rho = 0.26$  bit/s/Hz bits per channel use, which is larger than the maximum throughput achieved by all other schemes. On the other hand, for  $n\bar{W} = 2$  J/slot, the maximum throughput of IR-HARQ with and without subset selection is larger than that of Repetition-HARQ.

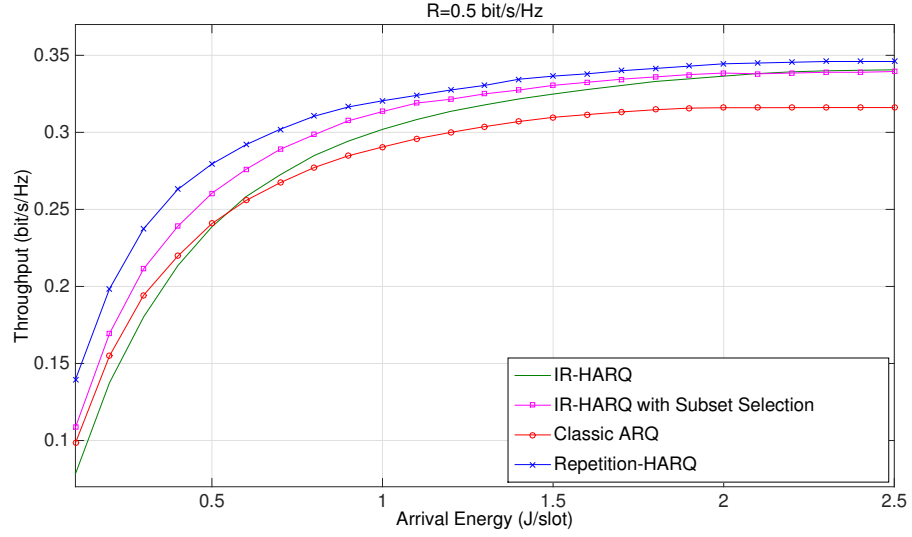


Figure 5.2: Throughput versus the arrival energy rate of  $n\bar{W}$  for the code rate of  $R = 0.5$  bits/s/Hz,  $n = 10^6$ .

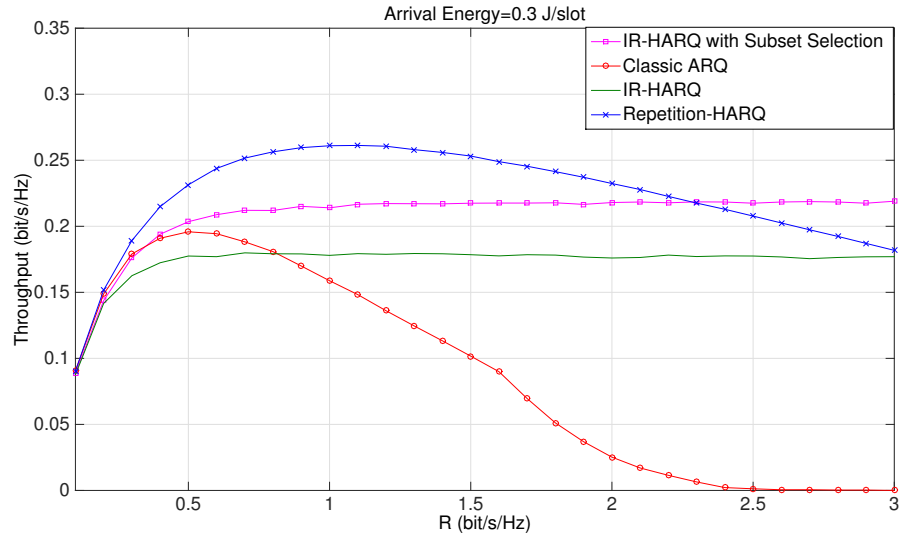


Figure 5.3: Throughput versus the code rate  $R$  for the energy arrival  $n\bar{W} = 0.3$  J/slot.



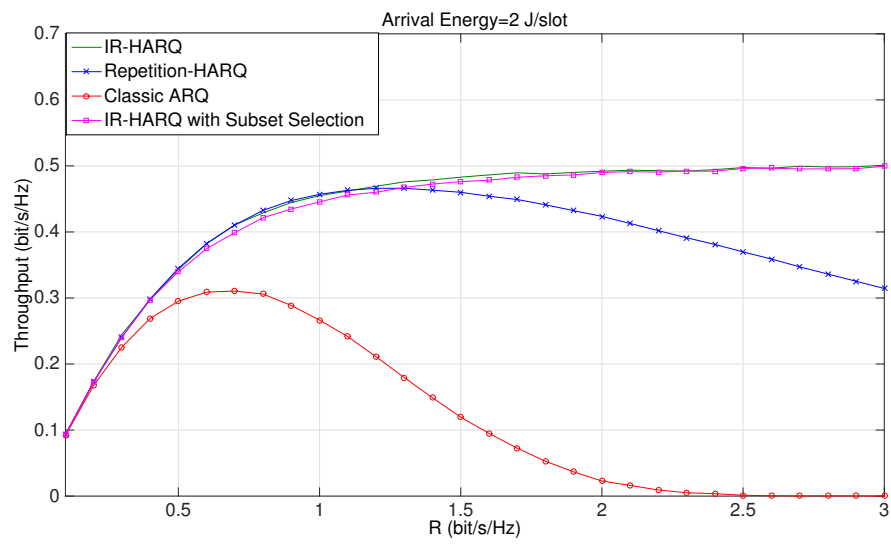


Figure 5.4: Throughput versus the code rate  $R$  for the energy arrival  $n\bar{W} = 2$  J/slot.

## Chapter 6

### Energy-Aware Downlink Scheduling in LTE-Advanced Networks

LTE cellular networks have become ubiquitous. The present deployments are predominantly based on Release 8, whereas newer ones (Release 13 is currently being finalized) will incorporate features that have been standardized in later releases. The key features of the LTE downlink (DL) are OFDMA and MIMO. As a result, a subframe scheduler that exploits these two features is a key component of an LTE base station. In subframe scheduling, which is done every millisecond, portions of the available bandwidth, referred to as Resource Blocks (RBs), are assigned to users. Data is sent to each scheduled user over its assigned RBs after an encoding process. The latter process comprises an outer encoder to generate up to two codewords and an inner encoder for transmit diversity or spatial multiplexing (spatial precoding.) A rich body of work has investigated the problem of subframe scheduling. Due to the fine timescale and processing power limitations, practical implementations greatly favor deterministic and low-complexity algorithms. Thus, in our view the contributions [57] and [58] have been particularly influential in demonstrating that the multi-carrier DL scheduling problem with finite queues and with two MIMO modes, respectively, can be approximately solved (with a constant-factor guarantee) by simple greedy algorithms. These works were then extended in [59] which imposed several practical LTE constraints. Recent advances include [60] and [61] which have incorporated and analyzed the impact of quantized (imperfect) channel state information (CSI) arising from specific LTE CSI feedback modes.

Our focus in this chapter is to incorporate energy efficiency in the LTE subframe scheduler design. The main motivation behind energy efficiency is to reliably transmit as

many bits as possible for every joule of energy spent, subject to QoS guarantees (cf. [62] for a comprehensive survey). Recent research in this area has also considered renewable energy sources [63,64]. A key aspect that arises when considering LTE networks is that certain signal attributes cannot (or should not) be changed dynamically at the time scales of subframe scheduling. These attributes for instance include:

- The choice of the transmitting node (or base station) especially in deployments with non-fibre backhaul.
- The choice of the transmit power level on an RB when it is assigned to a user.
- The maximum transmit rank (or number of data streams) that can be transmitted by a node.

These restrictions are a consequence of the limited control signaling support provided in the standard. Keeping in mind these limitations, we follow an approach for energy-aware subframe scheduler design that generalizes the one outlined in [63], which controlled the number of assigned RBs and the number of activated transmit antennas. The latter approach enabled [63] to incorporate renewable energy sources and is eminently suitable for energy efficiency as well. In our generalization, we start by adopting the transmission mode concept from [59] that can model the important LTE features and constraints, and also allows us to bound the maximum assignable transmit rank. The latter allows us to control the number of activated RF chains, which is proportional to the consumed baseband circuit power and subsumes the control of the number of activated physical antennas. Furthermore, we place a price on utilizing each RB and impose a cost constraint that a feasible set of scheduled RBs must satisfy, thereby significantly generalizing the control of number of allocated RBs advocated in [63]. Our framework allows us to approximately solve the single-cell subframe scheduling problem under a variety of important but seemingly intractable constraints that are discussed in the next section. Consequently, our proposed algorithm can be used as a sub-routine in energy efficient resource allocation, which has received significant attention spanning: point-to-point links [65,66], the single-cell downlink [67,68] followed

by the multi-cell downlink [68–70]. These works formulate resource allocation problems that are typically continuous optimization ones and fractional programming has emerged as a popular tool to solve such problems [62]. On the other hand, we rely on discrete combinatorial optimization tools which can directly consider resources (that are in general discrete) without resorting to any continuous relaxations.

The problem we formulate (not surprisingly) is NP hard, but unlike those in [57, 59], it cannot in general be reformulated to a form on which the classical greedy algorithm [71] yields a meaningful guarantee. This is a major complication since the designed algorithm must be simple and deterministic. Consequently, we turn to the multiplicative updates based method that was used to maximize a submodular set function subject to multiple knapsack constraints [72] (cf. the definitions in the appendix). Direct adaptations of the algorithms from [72] are either not possible or do not yield a useful guarantee since the problem at hand involves the combination of a matroid and (non-binary) knapsack constraints. We thus design a novel algorithm and build upon the intricate analysis developed in [72] to demonstrate that our algorithm achieves a meaningful guarantee. In particular, we show that it achieves an approximation factor that scales as  $1/\ln(n)$ , where  $n$  is proportional to the problem dimension. A key feature of our algorithm is that we work with a submodular form of the cost constraint instead of a linear constraint. This is a departure from the method in [73] where submodular constraints are successively approximated as linear. We are able to considerably enhance the analysis of [72] to incorporate such a submodular constraint and show that our approach offers significant performance improvements as well.

## 6.1 Problem formulation

### 6.1.1 Optimization framework

We consider the downlink sub-frame scheduling problem with the objective of maximizing the weighted sum of throughputs. We begin by discussing several practical constraints:

**[C1]:** An orthogonality constraint on each RB, i.e., at most one user from the

active user set  $\mathcal{U}$  can be scheduled on each RB. Also, the number of users scheduled in a subframe cannot exceed a certain threshold,  $\bar{K}$ . The latter constraint helps in reducing the control channel signaling overhead.

**[C2]:** The transmission mode on all the RBs allocated to the same user needs to be identical over a subframe. Each transmission mode must be selected from a given finite set  $\mathcal{M}$  of cardinality  $M$ . This generic constraint is very useful in incorporating several practical and mandatory ones. For instance, a mode can represent the transmit (spatial) rank assigned to the user (LTE Release 10 and beyond), or the choice between spatial multiplexing and transmit diversity (LTE Release 8 and beyond), or it can indicate a transmit precoder drawn from a finite codebook (LTE Release 8). Thus, by imposing the one transmission mode per scheduled user constraint, we can account for mandatory constraints imposed by several LTE releases.

**[C3]:** The bit loading on any RB must conform to choices permissible under the selected mode. This constraint facilitates in modeling the finite modulation and coding scheme (MCS) constraints. In addition, to address the bursty user traffic demands, the total number of bits assigned to each scheduled user  $u$  must not exceed its given queue (buffer) size  $Q_u$ . Note that the user traffic demands are typically bursty, hence it is very important to incorporate finite queue sizes.

**[C4]:** The set of all the RBs that are utilized must satisfy one or more linear cost constraints. Incorporating these cost constraints together with the maximum rank constraint, is very useful for energy-aware scheduler design.

Having described the constraints that we impose, we define the expected throughput obtained upon scheduling user  $u$  with mode  $m$  on RB  $n$  and loading  $b_n$  bits, as

$$b_n(1 - p_{u,n}^m(b_n)), \quad b_n \in \mathcal{B}^m, \quad (6.1)$$

where  $p_{u,n}^m(b_n)$  is the corresponding block error probability. The throughput in (6.1) is zero at  $b_n = 0$ .  $\mathcal{B}^m$  denotes the set of all possible bit loadings that can be done on any RB under mode  $m$  and we assume  $0 \in \mathcal{B}^m$ . The set  $\mathcal{B}^m$  can be any arbitrary finite set

or it can be an uncountable set. We define

$$B_{\max}^m = \max\{b : b \in \mathcal{B}^m\}.$$

Then, letting  $\alpha_u > 0$  denote the given (input) weight assigned to user  $u$  and maximizing over the bit loading, we obtain the maximum expected weighted throughput

$$r_{u,n}^m = \alpha_u \max_{b_n \in \mathcal{B}^m} \{b_n(1 - p_{u,n}^m(b_n))\}. \quad (6.2)$$

The throughput model in (6.1) allows us to accommodate imperfections (or errors) in the channel state information at the scheduler (cf. [60]) and is more general than the typical model used in many previous works (e.g. [57, 58]), where any variable number of bits  $b_n$  (subject to an upper bound) can be allocated on RB  $n$  and then all of them are then successfully delivered. Indeed, this latter model, henceforth referred to as the 0 – 1 throughput model, assumes that each error probability is a step function, i.e.,

$$p_{u,n}^m(b_n) = \mathbf{1}(b_n > B_{u,n}^m), \quad (6.3)$$

$B_{u,n}^m > 0$  denotes the maximal number of information bits that can be successfully loaded on RB  $n$  when scheduling user  $u$  with mode  $m$ . In addition, under this model the set  $\mathcal{B}^m$  is assumed to be an interval  $[0, B_{\max}^m]$ , with  $B_{\max}^m \geq B_{u,n}^m$ , for all  $u, m, n$ .

We define  $x_{u,n}^m$  as the indicator variable which is one if user  $u$  is scheduled on RB  $n$  with transmission mode  $m$ , and zero otherwise. Let  $K, M$  be the total number of users and transmission modes, respectively. Let  $\mathcal{N} = \{1, \dots, N\}$  denote the set of available RBs. We use the non-negative scalars  $\{a_n | n \in \mathcal{N}\}$  to denote the normalized prices associated with each of the RBs. Without loss of generality, we assume  $a_n \in [0, 1]$ , for all  $n \in \mathcal{N}$  and impose a constraint  $\sum_{n \in \mathcal{R}} a_n \leq 1$  that each feasible set  $\mathcal{R} \subseteq \mathcal{N}$  of allocated RBs must satisfy. Notice here that by setting  $a_n = 1/J$  for all  $n$  and some  $J \geq 1$ , we can ensure that no more than  $J$  RBs are chosen. Another example is where the cost constraint can impose a sum power constraint, wherein higher prices are assigned to some RBs compared to others. The former ones can be those RBs on

which transmission with a boosted power is permitted to service edge users in an FFR (Fractional Frequency Reuse) configuration. Furthermore, the bound on the maximal assignable transmit rank can be used to limit the number of RF chains that need to be activated. Then, for each instance, the input comprises of the set of prices, the maximum assignable rank, user weights and buffer sizes, the user limit and the throughput in (6.1) for all  $u, m, n$ . To improve readability some proofs for the claims for this chapter have been deferred to the appendix.

### 6.1.2 Backlogged traffic model

Under this model, using (6.2) we can formulate the scheduling problem as

$$\max_{\{x_{u,n}^m \in \{0,1\}\}} \sum_{u=1}^K \sum_{n=1}^N \sum_{m=1}^M r_{u,n}^m x_{u,n}^m \quad (6.4a)$$

$$\text{subject to } \sum_{m=1}^M \sum_{u=1}^K x_{u,n}^m \leq 1, \quad 1 \leq n \leq N, \quad (6.4b)$$

$$\sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq 1, \quad 1 \leq u \leq K, \quad (6.4c)$$

$$\sum_{u=1}^K \sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq \bar{K}, \quad (6.4d)$$

$$\sum_{n=1}^N \sum_{u=1}^K \sum_{m=1}^M a_n x_{u,n}^m \leq 1. \quad (6.4e)$$

Note that the set of constraints (6.4b) requires that on each RB at most one user with one transmission mode can be scheduled. Each of the  $M$  modes conforms to the given maximum rank bound. The set of constraints (6.4c) stipulates that each user can only have one transmission mode (i.e., one common transmission mode is used across all RBs allocated to the user). The constraint (6.4d) requires that the maximum number of scheduled users cannot exceed  $\bar{K}$ , for any given  $\bar{K}$  such that  $1 \leq \bar{K} \leq K$ . The last constraint (6.4e) dictates that the sum cost of all occupied RBs should be no greater than unity.

### 6.1.3 Finite queue model

Under this model, we assume that user  $u$  has a finite queue size  $Q_u$ . Denote  $b_{u,n}^m$  as the number of information bits allocated to user  $u$  on RB  $n$  with transmission mode  $m$  (these bits are however counted towards the weighted sum throughput objective only when user  $u$  is assigned RB  $n$  with mode  $m$ , i.e., only when  $x_{u,n}^m = 1$ ). The scheduling problem is then formulated as,

$$\max_{\{x_{u,n}^m, b_{u,n}^m\}} \sum_{u=1}^K \sum_{n=1}^N \sum_{m=1}^M \alpha_u x_{u,n}^m b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m)) \quad (6.5a)$$

$$\text{subject to } \sum_{m=1}^M \sum_{u=1}^K x_{u,n}^m \leq 1, \quad 1 \leq n \leq N, \quad (6.5b)$$

$$\sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq 1, \quad 1 \leq u \leq K, \quad (6.5c)$$

$$\sum_{u=1}^K \sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq \bar{K}, \quad (6.5d)$$

$$\sum_{n=1}^N \sum_{u=1}^K \sum_{m=1}^M a_n x_{u,n}^m \leq 1, \quad (6.5e)$$

$$\sum_{m=1}^M \sum_{n=1}^N b_{u,n}^m \leq Q_u, \quad 1 \leq u \leq K, \quad (6.5f)$$

$$x_{u,n}^m \in \{0, 1\} \text{ \& } b_{u,n}^m \in \mathcal{B}^m, \forall u, m, n, \quad (6.5g)$$

where the set of constraints (6.5f) model the queue size limit for each user.

### 6.1.4 Hardness result

To prove the hardness of the problems in (6.4) and (6.5) it suffices to consider the 0 – 1 throughput model (6.3) and instances where the sum cost constraint is irrelevant (i.e.,  $\sum_{n \in \mathcal{N}} a_n \leq 1$ ). We then see that (6.4) subsumes the problem formulated in [58] since it considers more than two transmission modes. On the other hand, (6.5) subsumes the problem formulated in [57] since it considers more than one mode. Invoking the hardness results established in [57, 58], we can deduce the following result.

**Remark 6.1** *Problem (6.4) is NP-hard. There exists a  $\delta > 0$ , such that it is NP-hard*



to obtain  $(1 - \delta)$ -approximation to problem (6.5).

The simplified version of the problem in (6.5) considered in [57] is a combinatorial auction problem without any cost constraint and with submodular valuations, i.e., where the RBs have to be distributed among the users with each user's utility being a submodular set function. For this problem the classical greedy algorithm [71] is well suited. Regarding the original version in (6.5), we can show the following.

**Theorem 6.1** *For the 0 – 1 throughput model, assuming  $\bar{K} \geq K$ , problem (6.5) is a combinatorial auction problem with a linear cost constraint and valuations that are fractionally sub-additive (but not necessarily submodular).*

The proof is found in the appendix.

## 6.2 A Unified Scheduling Algorithm

We develop a unified algorithm to solve the downlink scheduling problems in (6.4) and (6.5). We show that this algorithm achieves a good approximation ratio under both the traffic models.

Our objective is to construct a set  $\mathcal{S}$  of scheduled users, a set  $\mathcal{R}$  of selected RBs with  $\sum_{n \in \mathcal{R}} a_n \leq 1$ , a transmission mode and a distinct set of RBs in  $\mathcal{R}$  allocated to each user  $u \in \mathcal{S}$ . To achieve this objective, we adopt a two staged approach where the first stage selects a set of users in which each user is assigned one mode and a distinct set of RBs. This selection is feasible with respect to all constraints of (6.4) and (6.5) except for the cost constraint. Consequently, in the second stage the output of the first stage is pruned in order to obtain a (fully) feasible solution. The key insight we use is to enforce the cost constraint in a *soft* manner in the first stage. This allows us to obtain a solution from the first stage that is close to optimal, while not significantly violating the cost constraint. Then, pruning the solution so obtained yields a feasible solution for which a good approximation guarantee can be derived. We remark here that we select the cost constraint as the one which will be enforced in a soft manner in the first stage. This is because after appropriate reformulation (as will be revealed later), all the other constraints together form a single matroid and can be strictly enforced in the

first stage while ensuring an approximate optimality. The pseudo-code of our two-stage algorithm is illustrated in Algorithm 1, which we next proceed to elucidate.

We begin by considering the first stage of the algorithm. We denote  $\mathcal{U}$  as the set of candidate users that have not been selected so far, and we initialize  $\mathcal{U} = \{1, \dots, K\}$  together with  $\mathcal{S} = \emptyset$ . We define a current value  $V_n$  for each RB  $n \in \mathcal{N}$ , representing the weighted throughput of the current allocation on that RB, and initially we set  $V_n = 0$  for all  $n \in \mathcal{N}$ . Note that we can also consider each  $V_n$  to be the weighted throughput barrier that a new selection (of user and mode) should exceed on RB  $n$  in order to obtain an assignment. After each iteration we use the scalar  $w$  to track the accumulated cost of all the occupied RBs, i.e.,

$$w = \lambda^{\sum_{n \in \mathcal{N}} a_n \mathbf{1}(V_n > 0)}, \quad (6.6)$$

and  $w$  is initialized to one. Further, we let  $\lambda$  be a scalar that we can initialize to any value strictly greater than 1. The role of  $\lambda$  is to enforce the cost constraint in a soft manner. The approximation ratio does depend on  $\lambda$  and its impact is investigated in the simulation results.

The first stage of our algorithm is iterative. At each iteration, in order to select a new user (and its mode), we need to solve the following problem for each candidate user  $u \in \mathcal{U}$  and each mode  $m \in \{1, \dots, M\}$ :

$$\begin{aligned} & \max_{\{v(u,n,m)\}} \sum_{n=1}^N \max\{v(u,n,m) - V_n, 0\} \\ & \text{subject to } \sum_{n=1}^N a_n \mathbf{1}(V_n = 0 \ \& \ v(n,u,m) > 0) \leq \lambda/w, \\ & \sum_{n=1}^N a_n \mathbf{1}(v(n,u,m) > 0) \leq 1, \end{aligned} \quad (6.7)$$

where  $v(u,n,m)$  denotes the weighted throughput obtained upon assigning user  $u$  to RB  $n$  with mode  $m$ . The first constraint in (6.7) ensures that the sum cost over all assigned RBs that were hitherto unoccupied; i.e., each assigned RB  $n$  for which  $V_n = 0$ , does not exceed  $\lambda/w$ . Imposing this constraint at each iteration ensures that upon

termination of the first stage, the sum cost over all occupied RBs is not too large but can exceed unity (soft constraint). The second constraint ensures that the sum cost of all assigned RBs does not exceed unity. Let  $\{\hat{v}(u, n, m)\}_{n \in \mathcal{N}}$  denote a good feasible solution to (6.7). Then, we define

$$g(u, m) = \sum_{n=1}^N \max\{\hat{v}(u, n, m) - V_n, 0\},$$

to be the net gain of adding user  $u$  to  $\mathcal{S}$  with transmission mode  $m$ . The exact computation of  $\hat{v}(u, n, m)$  (and thus of  $g(u, m)$ ) depends on the traffic model and will be discussed shortly. Next, let

$$(u^*, m^*) = \operatorname{argmax}_{u \in \mathcal{U}, m \leq M} g(u, m).$$

We then add the user  $u^*$  to the scheduled user list (and remove it from  $\mathcal{U}$ ), assign its transmission mode to be  $m^*$ . For each RB  $n \in \mathcal{N}$ , if  $\hat{v}(u^*, n, m^*) > V_n$ , then it is allocated to user  $u^*$  and removed from its previous allocation if it was allocated previously to any user. This ensures that each user is selected at most once (with one mode) and that each RB is always assigned to at-most one user. Note that the RB allocation is subject to change as an assigned RB may be re-allocated to another user later or that RB itself may be dropped from the set  $\mathcal{R}$  that is eventually selected.  $V_n$  is now updated as

$$V_n = \max\{V_n, \hat{v}(u^*, n, m^*)\}, \quad n \in \mathcal{N}, \quad (6.8)$$

followed by updating  $w$  according to (6.6). The first stage of the algorithm terminates when any of the following conditions are satisfied:

1.  $\mathcal{U} = \emptyset$ ,
2.  $|\mathcal{S}| = \bar{K}$ ,
3.  $g(u^*, m^*) = 0$ .

In the second stage we first consider the set of all selected RBs  $\mathcal{R}' = \{n : V_n > 0\}$ .

We declare the obtained solution feasible if  $\sum_{n \in \mathcal{R}'} a_n \leq 1$  and set  $\mathcal{R} = \mathcal{R}'$ . Otherwise, we follow a *pruning* procedure detailed below. Under this procedure, we solve the following standard knapsack problem

$$\begin{aligned} & \max_{\mathcal{R} \subseteq \mathcal{R}'} \left\{ \sum_{n \in \mathcal{R}} V_n \right\} \\ & \text{s.t. } \sum_{n \in \mathcal{R}} a_n \leq 1, \end{aligned} \quad (6.9)$$

using any suitable approximation algorithm [74] to determine the feasible subset  $\mathcal{R}$ . Finally, considering each user  $u$  which is assigned one or more RBs within the set  $\mathcal{R}$ , we re-select its best mode without changing the RB allocation. Note that the latter optimization is decoupled across users and is simple since the RB allocations are not altered.

As promised above, we now consider the sub-problem (6.7) that needs to be solved at each iteration, for each one of the two following traffic models.

**Backlogged traffic model:** Under the backlogged traffic model, without loss of optimality, we can deduce that if the  $n^{\text{th}}$  RB is selected then  $v(u, n, m)$  must be the maximum weighted expected throughput of allocating user  $u$  on RB  $n$  with transmission mode  $m$ , i.e.,

$$v(u, n, m) = z_n r_{u,n}^m, \quad (6.10)$$

where  $z_n \in \{0, 1\}$  is an indicator variable that is one if RB  $n$  is chosen and is zero otherwise. Hence,  $g(u, m)$  can be calculated by solving the following problem

$$\begin{aligned} & \max_{\{z_n \in \{0,1\}, \forall n\}} \sum_{n=1}^N \max\{z_n r_{u,n}^m - V_n, 0\} \\ & \text{s.t. } \sum_{n=1}^N a_n \mathbf{1}(V_n = 0) z_n \leq \lambda/w, \\ & \sum_{n=1}^N a_n z_n \leq 1. \end{aligned} \quad (6.11)$$

**Finite queue model:** Given  $u, m$  and the queue size  $Q_u$  for user  $u$ , the weighted

---

**Algorithm 1** Unified Scheduling Algorithm
 

---

```

1:  $\mathcal{U} = \{1, \dots, K\}, \mathcal{S} = \emptyset, V_n = 0$  for  $1 \leq n \leq N, \lambda > 1, w = 0$ .
2: while Termination conditions not satisfied do
3:   for all  $u \in \mathcal{U}$  do
4:     for  $m = 1$  to  $M$  do
5:       Compute the value  $\hat{v}(u, n, m)$  on each RB  $n$  and the net gain
          $g(u, m)$ , while satisfying  $\sum_{n=1}^N a_n \mathbf{1}(V_n = 0 \ \& \ \hat{v}(n, u, m) > 0) \leq \lambda/w$  and
          $\sum_{n=1}^N a_n \mathbf{1}(\hat{v}(n, u, m) > 0) \leq 1$ 
6:     end for
7:   end for
8:    $(u^*, m^*) = \arg \max_{u \in \mathcal{U}, 1 \leq m \leq M} g(u, m)$ .
9:   If  $g(u^*, m^*) = 0$ , go to 19.
10:  Add user  $u^*$  to  $\mathcal{S}$  with mode  $m^*$  and remove it from  $\mathcal{U}$ .
11:  for  $n = 1, \dots, N$  do
12:    if  $\hat{v}(u^*, n, m^*) > V_n$  then
13:      Allocate RB  $n$  to user  $u^*$ , after removing it from any previously assigned
        user.
14:       $V_n = \hat{v}(u^*, n, m^*)$ .
15:    end if
16:  end for
17:  Update  $w = \lambda \sum_{n=1}^N a_n \mathbf{1}(V_n > 0)$ 
18: end while
19: Determine  $\mathcal{R}' = \{n : V_n > 0\}$ 
20: if  $\sum_{n \in \mathcal{R}'} a_n > 1$  then
21:   Determine a subset  $\mathcal{R} \subseteq \mathcal{R}'$  such that  $\sum_{n \in \mathcal{R}} a_n \leq 1$  via pruning.
22: else
23:   Set  $\mathcal{R} = \mathcal{R}'$ 
24: end if
25: Re-select the optimal transmission mode for each user over its assigned RBs within
    the set  $\mathcal{R}$  of selected RBs.

```

---

throughput  $v(u, n, m)$  is a function of the bit loading  $b_n$  on any selected RB  $n$ , i.e.,  $v(u, n, m) = z_n \alpha_u b_n (1 - p_{u,n}^m(b_n))$ . As a result, the major issue here is to find a good bit allocation subject to the queue size constraint. Towards that end, we define a set  $\mathcal{B}_u^m = \{b \in \mathcal{B}^m : b \leq Q_u\}$ . Without loss of generality, we assume that the set  $\mathcal{B}_u^m$  has at-least one strictly positive entry (otherwise we can simply remove mode  $m$  as a candidate mode for user  $u$ ). Then, the problem of interest can be posed as the following.

$$\begin{aligned}
& \max \sum_{n=1}^N \max\{z_n \alpha_u b_n (1 - p_{u,n}^m(b_n)) - V_n, 0\} \\
& \text{s.t. } b_n \in \mathcal{B}_u^m, \forall n; z_n \in \{0, 1\}, \\
& \sum_{n=1}^N b_n \leq Q_u, \\
& \sum_{n=1}^N a_n \mathbf{1}(V_n = 0) z_n \leq \lambda/w, \\
& \sum_{n=1}^N a_n z_n \leq 1.
\end{aligned} \tag{6.12}$$

Clearly, the problem in (6.12) subsumes that in (6.11). Notice also that the penultimate constraint in (6.12) is vacuous (in lieu of the last one) whenever  $\lambda/w \geq 1$ .

We next propose a simple approximation algorithm to solve both (6.11) and (6.12). We let

$$\bar{r}_{u,n}^m = \alpha_u \max_{b_n \in \mathcal{B}_u^m} \{b_n (1 - p_{u,n}^m(b_n))\},$$

with

$$\bar{B}_{u,n}^m = \arg \max_{b_n \in \mathcal{B}_u^m} \{b_n (1 - p_{u,n}^m(b_n))\},$$

and define a subset of candidate RBs

$$\mathcal{R}_{\text{cand}} = \{n \in \mathcal{N} : \bar{r}_{u,n}^m > V_n \text{ \& } a_n \mathbf{1}(V_n = 0) \leq \lambda/w\}.$$

Clearly there is no reason to allocate to user  $u$  any RBs outside of  $\mathcal{R}_{\text{cand}}$ . Moreover, if  $\mathcal{R}_{\text{cand}}$  is empty we can simply set  $\hat{v}(u, n, m) = 0$  for all  $n \in \mathcal{N}$ . Otherwise, we define

three conditions

$$\begin{aligned}
\mathbf{D1} & : \sum_{n \in \mathcal{R}_{\text{cand}}} \bar{B}_{u,n}^m \leq Q_u, \\
\mathbf{D2} & : \sum_{n \in \mathcal{R}_{\text{cand}}} a_n \mathbf{1}(V_n = 0) \leq \lambda/w, \\
\mathbf{D3} & : \sum_{n \in \mathcal{R}_{\text{cand}}} a_n \leq 1.
\end{aligned} \tag{6.13}$$

If all these conditions are satisfied (i.e., are true) we can optimally solve (6.12) by setting

$$\hat{v}(u, n, m) = \begin{cases} \bar{r}_{u,n}^m & n \in \mathcal{R}_{\text{cand}} \\ 0 & n \in \mathcal{N} \setminus \mathcal{R}_{\text{cand}} \end{cases}$$

For the remaining possibilities we first determine the effective normalized RB costs

$$\bar{a}_n = a_n \mathbf{1}(V_n = 0) / (\lambda/w), \forall n.$$

Then, we set scalar  $\zeta_1 = 1$  ( $\zeta_3 = 1$ ) if the condition D1 (D3) is not satisfied and set  $\zeta_1 = 0$  ( $\zeta_3 = 0$ ) otherwise. Another scalar is defined as  $\zeta_2 = 0$  whenever the condition D2 is satisfied or when  $\lambda/w \geq 1$  (which is always true in the first iteration) and is set to be one otherwise. Finally, we select any scalar  $\theta$  such that  $\theta > \zeta_1 + \zeta_2 + \zeta_3$ . Then, we determine a good feasible (suboptimal) solution to (6.12) using the greedy procedure that is outlined in Algorithm 2. Some comments on Step 4 of Algorithm 2 are in order. We suppose that the locally optimal RB  $\hat{\ell}$  together with the corresponding optimal bit loading  $\hat{b}_{\hat{\ell}}$  (which we assume must exist) can be efficiently determined. This is indeed true when the set  $\mathcal{B}^m$  (and thus  $\mathcal{B}_u^m$ ) is finite. It is also true for the 0 – 1 throughput model in which an optimal bit loading for the inner maximization on any RB  $\ell$  is given by  $\hat{b}_{\ell} = \bar{B}_{u,\ell}^m$ . Further, in case that the optimal bit loading on any RB is not unique, we select the largest one among all such optimal loadings. While the approximation guarantee that will be proved next holds true for any arbitrarily fixed tie breaking rule, we have seen in our simulation results that the recommended method works better.

We now proceed to establish an approximation guarantee for Algorithm 1. Towards

---

**Algorithm 2** Greedy Bit Allocation
 

---

- 1: Set  $\hat{v}(u, n, m) = 0$ ,  $\forall n = 1, \dots, N$  and  $G = 0$  and  $\check{V}_n = V_n$ ,  $\forall n \in \mathcal{R}_{\text{cand}}$
  - 2: Initialize  $\theta, \zeta_1, \zeta_2, \zeta_3$  and  $\bar{a}_n$ ,  $\forall n$ .
  - 3: **while**  $\zeta_1 + \zeta_2 + \zeta_3 \leq \theta$  **do**
  - 4:   Determine  
        $\hat{\ell} = \arg \max_{\ell \in \mathcal{R}_{\text{cand}}} \max_{\substack{b_\ell \in \mathcal{B}_u^m \\ b_\ell > 0}} \left\{ \frac{\alpha_u b_\ell (1 - p_{u,\ell}^m(b_\ell)) - \check{V}_\ell}{\zeta_1 b_\ell / Q_u + \zeta_2 \bar{a}_\ell + \zeta_3 a_\ell} \right\}$  and let  $\hat{b}_{\hat{\ell}}$  denote the corresponding optimal bit loading.
  - 5:   **if**  $\zeta_1 \theta^{\hat{b}_{\hat{\ell}}/Q_u} \leq \theta$  and  $\zeta_2 \theta^{\bar{a}_{\hat{\ell}}} \leq \theta$  and  $\zeta_3 \theta^{a_{\hat{\ell}}} \leq \theta$  **then**
  - 6:     Set  $\hat{v}(u, \hat{\ell}, m) = \alpha_u \hat{b}_{\hat{\ell}} (1 - p_{u,\hat{\ell}}^m(\hat{b}_{\hat{\ell}}))$  and  $G = G + \hat{v}(u, \hat{\ell}, m) - \check{V}_{\hat{\ell}}$
  - 7:     Set  $\check{V}_{\hat{\ell}} = \hat{v}(u, \hat{\ell}, m)$
  - 8:   **end if**
  - 9:   Update  $\zeta_1 = \zeta_1 \theta^{\hat{b}_{\hat{\ell}}/Q_u}$  and  $\zeta_2 = \zeta_2 \theta^{\bar{a}_{\hat{\ell}}}$  and  $\zeta_3 = \zeta_3 \theta^{a_{\hat{\ell}}}$
  - 10: **end while**
  - 11: Determine  $\hat{\ell} = \arg \max_{\ell \in \mathcal{R}_{\text{cand}}} \{\bar{r}_{u,\ell}^m - V_\ell\}$
  - 12: **if**  $\bar{r}_{u,\hat{\ell}}^m - V_{\hat{\ell}} > G$  **then**
  - 13:   Set  $\hat{v}(u, \hat{\ell}, m) = \bar{r}_{u,\hat{\ell}}^m$  and  $\hat{v}(u, n, m) = 0$ ,  $\forall n \neq \hat{\ell}$ .
  - 14: **end if**
- 

this end, we first derive a guarantee for Algorithm 2, where we need to consider only the case with non-empty  $\mathcal{R}_{\text{cand}}$  and  $\zeta_1 + \zeta_2 + \zeta_3 > 0$ .

**Proposition 6.1** *Algorithm 2 is a constant-factor approximation algorithm for the problem in (6.12), where the constant factor is given by  $\eta = \frac{1}{2} \left( 1 - \frac{1}{\exp(1/\theta)} \right)$ .*

**Proof:** We recall the definition of the set  $\mathcal{R}_{\text{cand}}$  and define a ground set of tuples  $\tilde{\Psi} = \{(\ell, b_\ell), \forall \ell \in \mathcal{R}_{\text{cand}} \ \& \ b_\ell \in \mathcal{B}_u^m : b_\ell > 0\}$ . Note that for the same RB  $\ell$  we can have several distinct tuples in  $\tilde{\Psi}$  corresponding to different bit loading  $b_\ell$ . Then, we define a normalized non-negative set function  $f : 2^{\tilde{\Psi}} \rightarrow \mathbb{R}_+$ , such that  $f(\emptyset) = 0$  and for all other subsets  $\mathcal{A} \subseteq \tilde{\Psi}$ , we have

$$f(\mathcal{A}) = \sum_{\ell' \in \mathcal{R}_{\text{cand}}} \max_{(\ell, b_\ell) \in \mathcal{A}} \{ \mathbf{1}(\ell' = \ell) [\alpha_u b_\ell (1 - p_{u,\ell}^m(b_\ell)) - V_\ell]^+ \}.$$



We can now express the problem in (6.12) as

$$\begin{aligned}
& \max_{\mathcal{A} \subseteq \tilde{\Psi}} \{f(\mathcal{A})\} \\
& \text{s.t.} \quad \sum_{(\ell, b_\ell) \in \mathcal{A}} b_\ell \leq Q_u, \\
& \quad \sum_{(\ell, b_\ell) \in \mathcal{A}} a_\ell \mathbf{1}(V_\ell = 0) \leq \lambda/w, \\
& \quad \sum_{(\ell, b_\ell) \in \mathcal{A}} a_\ell \leq 1.
\end{aligned} \tag{6.14}$$

It can be readily verified that the set function  $f(\cdot)$  is a normalized non-decreasing submodular set function (cf. the definitions given in the appendix). Consequently, the problem in (6.14) is that of maximizing such a set function subject to three linear packing (knapsack) constraints. The conditions in (6.13) identify which of the three knapsack constraints are vacuous. For each constraint  $i$  the corresponding scalar  $\zeta_i$  is defined to be one only if it is relevant. Then, in case only one of the constraints is relevant, Algorithm 2 reduces to the greedy algorithm of [75], which considered submodular set function maximization subject to one knapsack constraint, and yields a guarantee of  $1 - \frac{1}{\sqrt{e}}$ . Moreover, in each case, Algorithm 2 is a simple enhancement of the multiplicative updates based algorithm of [72], which considered submodular set function maximization subject to multiple general knapsack constraints. The direct adaptation of the algorithm from [72] would skip the check in Step 5. Then, if the solution obtained upon termination (of the while-loop) is infeasible, it would compare objective value attained by just the last tuple added, against that yielded by all the other selected tuples (i.e., the first to the penultimate one) together, and pick the choice yielding the larger objective value. Our enhancement, which allows us to unify the algorithms from [75] and [72], can be readily seen to yield a solution at-least as good as this direct adaptation. Consequently, we can claim the guarantee of  $\frac{1}{2} \left(1 - \frac{1}{\exp(1/\theta)}\right)$  for all  $\theta > 3$  established in [72] (for the problem with three knapsack constraints). The desired result then corresponds to the latter smaller guarantee.  $\square$

We are now ready to establish the approximation guarantee for Algorithm 1. Towards

that end, we define a ground set containing all possible 3-tuples  $\Psi = \{(u, m, \mathbf{b})\}$ , where for each such 3-tuple (or element)  $\underline{e} = (u, m, \mathbf{b})$ ,  $u$  denotes the user and  $m$  denotes the mode and  $\mathbf{b} = [b_1, \dots, b_N]^T \in \mathbb{R}_+^N$  is an  $N$  length bit loading vector such that for each RB  $\ell \in \mathcal{N}$ ,  $b_\ell \in \mathcal{B}^m$  and  $\sum_{\ell \in \mathcal{N}} b_\ell \leq Q_u$ . Further,  $\sum_{\ell \in \mathcal{N}} a_n \mathbf{1}(b_\ell > 0) \leq 1$ . With these qualifications we can conclude that  $\Psi$  includes all possible elements, where each such element represents a valid assignment of mode, RBs and bits to a user. Furthermore, we define a family of subsets of  $\Psi$ , denoted by  $\underline{\mathcal{I}}$ , as follows.

$$\underline{\mathcal{I}} = \left\{ \mathcal{A} \subseteq \Psi : \sum_{(u', m', \mathbf{b}') \in \mathcal{A}} \mathbf{1}(u' = u) \leq 1 \ \forall u \in \mathcal{U} \ \& \ |\mathcal{A}| \leq \bar{K} \right\}. \quad (6.15)$$

In words, any subset of 3-tuples from  $\Psi$  in which each user appears at-most once and whose cardinality does not exceed the user limit  $\bar{K}$  is a member of  $\underline{\mathcal{I}}$  and vice versa. Then, we define a normalized non-negative set function  $h : 2^\Psi \rightarrow \mathbb{R}_+$ , such that  $h(\emptyset) = 0$  and for all other subsets  $\mathcal{A} \subseteq \Psi$

$$h(\mathcal{A}) = \sum_{n \in \mathcal{N}} \max_{(u, m, \mathbf{b}) \in \mathcal{A}} \{\alpha_u b_n (1 - p_{u,n}^m(b_n))\}. \quad (6.16)$$

The following result follow from the basic definitions provided in the appendix.

**Lemma 6.1** *The set function  $h(\cdot)$  is a normalized non-decreasing submodular set function and  $(\Psi, \underline{\mathcal{I}})$  is a matroid.*

We can now reformulate the problem in (6.5) as

$$\begin{aligned} & \max_{\mathcal{A} \subseteq \Psi} \{h(\mathcal{A})\} \\ & \text{s.t. } \mathcal{A} \in \underline{\mathcal{I}} \\ & \sum_{(u, m, \mathbf{b}) \in \mathcal{A}} \sum_{n \in \mathcal{N}} a_n \mathbf{1}(b_n > 0) \leq 1. \end{aligned} \quad (6.17)$$

Some comments on this reformulation are in order. First, from the definition of  $\underline{\mathcal{I}}$  it follows that by restricting  $\mathcal{A} \in \underline{\mathcal{I}}$  we have ensured that each user is selected at-most once with one mode and that the associated bit loading is feasible, thereby meeting

the per-user mode and bit loading constraints of (6.5). Then, the definition of the set function  $h(\cdot)$  ensures that each RB is implicitly assigned to at-most one user (since only the weighted throughput of at-most one user is chosen via the  $\max(\cdot)$  function). Consequently, any feasible solution to (6.17) maps to a feasible one for (6.5) yielding the same objective value and vice versa. Thus, we have reformulated the problem in (6.5) as that of maximizing a submodular objective subject to one matroid and one knapsack constraint. A subtle and useful point is that, without loss of optimality, we can replace the linear knapsack constraint in (6.17) with a submodular one, i.e.,

$$\begin{aligned} & \max_{\mathcal{A} \subseteq \Psi} \{h(\mathcal{A})\} \\ \text{s.t. } & \mathcal{A} \in \underline{\mathcal{I}}, \\ & \sum_{n \in \mathcal{N}} a_n \mathbf{1}(\exists(u, m, \mathbf{b}) \in \mathcal{A} : b_n > 0) \leq 1. \end{aligned} \quad (6.18)$$

To see that the formulations in (6.17) and (6.18) are equivalent, we first note that any optimal solution for (6.17) is clearly feasible for (6.18). On the other hand, given any optimal solution  $\mathcal{O}^{\text{opt}}$  for (6.18), we can transform it to one that yields the same objective value but is feasible for (6.17), by considering each element  $\hat{\underline{e}} = (\hat{u}, \hat{m}, \hat{\mathbf{b}}) \in \mathcal{O}^{\text{opt}}$  and each RB  $n \in \mathcal{N}$ , and forcing  $\hat{b}_n = 0$  if  $\alpha_{\hat{u}} \hat{b}_n (1 - p_{\hat{u}, n}^{\hat{m}}(\hat{b}_n))$  is not the maximum weighted throughput on RB  $n$  among all elements in  $\mathcal{O}^{\text{opt}}$  (ties can be broken arbitrarily). This reformulation (6.18) with the submodular form of the constraint is used in designing Algorithm 1. The utility of (6.18) is that it can help sub-optimal algorithms since it expands the set of feasible solutions. Considering the problem in (6.18), let  $\mathcal{O}^{\text{opt}}$  denote an optimal solution and let  $\hat{\mathcal{O}}$  denote the final output yielded by Algorithm 1. Further, let  $\mathcal{O}$  denote the intermediate solution obtained after stage one of Algorithm 1. With some abuse of notation, let  $\mathbf{a}_{\mathcal{O}}$  denote total cost of all the RBs occupied under  $\mathcal{O}$ , i.e.,  $\mathbf{a}_{\mathcal{O}} = \sum_{n \in \mathcal{R}'} a_n$ , where  $\mathcal{R}'$  is defined in Step 19 of Algorithm 1. Recall that the intermediate solution  $\mathcal{O}$  need not be feasible with respect to the cost constraint, i.e., we can have  $\mathbf{a}_{\mathcal{O}} > 1$ . We next show that the intermediate solution  $\mathcal{O}$  at-most incurs a constant-factor loss with respect to the optimal one. To prove this result we build upon the brilliant proof of Theorem 1.3 of [72]. The key challenges we have to surmount are

that: the ground set  $\mathcal{B}^m$  can be uncountable and that the sub-problem (6.12) can only be approximately solved and that the submodular form of the constraint is employed. For convenience, we use the notation  $\mathbf{1}(n \in \mathcal{A})$  for any RB  $n \in \mathcal{N}$  and subset  $\mathcal{A} \subseteq \Psi$  to return one if there exists an element  $\underline{e} = (u, m, \mathbf{b}) \in \mathcal{A}$  under which the bit loading on the  $n^{\text{th}}$  RB is strictly positive ( $b_n > 0$ ), and return zero otherwise. Define

$$\mathbf{a}_{\mathcal{A}} = \sum_{n \in \mathcal{N}} a_n \mathbf{1}(n \in \mathcal{A}), \forall \mathcal{A} \subseteq \Psi.$$

Further, for any set  $\mathcal{A} \subseteq \Psi$  and element  $\underline{e} \in \Psi$ , we let

$$h_{\mathcal{A}}(\underline{e}) = h(\mathcal{A} \cup \underline{e}) - h(\mathcal{A}).$$

Then, let  $\underline{e}_j \in \Psi$ ,  $j = 1, \dots, J$  be the element (3-tuple) chosen by Algorithm 1 at the  $j^{\text{th}}$  iteration of stage-1 and  $J$  is the number of iterations after which stage-1 of the algorithm terminates. Letting  $w_j$  denote the scalar  $w$  after the  $j^{\text{th}}$  iteration of stage-1 of Algorithm 1, we see that  $w_j = \lambda^{\sum_{n=1}^N a_n \mathbf{1}(n \in \{\underline{e}_1, \dots, \underline{e}_j\})}$ ,  $\forall j = 1, \dots, J$ .

**Proposition 6.2** *The intermediate solution yielded by Algorithm 1,  $\mathcal{O}$ , satisfies,*

$$h(\mathcal{O}) \geq \frac{h(\mathcal{O}^{\text{opt}})\eta}{\eta + \lambda + 1}, \quad (6.19)$$

where  $\eta$  is the constant factor for the approximation guarantee of Algorithm 2 and  $\lambda$  is the parameter used for tracking the (soft) cost constraint as in (6.6).

**Proof:** We begin by expanding the elements of the ground set  $\Psi$  as the following ordered sequence of elements:  $\underline{\mathcal{V}} = \{\underline{e}_1, \mathcal{F}_1, \underline{e}_2, \mathcal{F}_2, \dots, \underline{e}_J, \mathcal{F}_J\}$ . Here,  $\mathcal{F}_1$  is the set formed by elements from  $\Psi \setminus \underline{e}_1$  that are also distinct from  $\{\underline{e}_j\}_{j=2}^J$ . In particular, all elements that share the same user as in  $\underline{e}_1$  belong to  $\mathcal{F}_1$  (such elements cannot be selected by the Algorithm at any later stage). In addition, each element,  $\underline{e}$ , that violates the soft cost constraint, i.e.,  $(\mathbf{a}_{\{\underline{e}_1, \underline{e}\}} - \mathbf{a}_{\underline{e}_1})w_1 > \lambda$  and does not belong to the

set  $\{\underline{e}_2, \dots, \underline{e}_J\}$ , is included in  $\mathcal{F}_1$ . Similarly, each set

$$\mathcal{F}_j \subseteq \Psi \setminus (\mathcal{F}_1 \cup \dots \cup \mathcal{F}_{j-1} \cup \{\underline{e}_1, \dots, \underline{e}_j\}), \quad j = 2, \dots, J-1$$

includes all elements that share the same user as in  $\underline{e}_j$ . In addition, each  $\mathcal{F}_j$  includes all elements that violate the soft cost constraint, i.e.,  $(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_j, \underline{e}\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_j\}})w_j > \lambda$  and do not belong to the set  $\mathcal{F}_1 \cup \mathcal{F}_2 \cup \dots \cup \mathcal{F}_{j-1} \cup \{\underline{e}_{j+1}, \dots, \underline{e}_J\}$ . The last set  $\mathcal{F}_J$  is the set of remaining elements  $\Psi \setminus (\mathcal{F}_1 \cup \dots \cup \mathcal{F}_{J-1} \cup \{\underline{e}_1, \dots, \underline{e}_J\})$ . Notice that an element  $\underline{e}$  in this set  $\mathcal{F}_J$  can have a user distinct from the one in  $\underline{e}_J$  and satisfy  $(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_J, \underline{e}\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_J\}})w_J \leq \lambda$ , only if either  $|\mathcal{O}| = J = \bar{K}$  or  $h_{\mathcal{O}}(\underline{e}) = 0$ . In other words, only when the algorithm terminated either due to the user limit being reached or no remaining element offering a strictly positive gain. Let  $\tilde{\mathcal{F}}_j = \underline{e}_j \cup \mathcal{F}_j$   $j = 1, \dots, J$ . The main intuition that we leverage from [72] is as follows. Consider a traversal of the ordered sequence  $\underline{\mathcal{V}}$  (starting from  $\underline{e}_1$ ) such that Algorithm 1 selects elements  $\underline{e}_1, \dots, \underline{e}_J$  while discarding the rest. Suppose we can show that at every instance in this traversal the cardinality of the set selected by Algorithm 1 remains at least a certain constant fraction of that selected by the optimal one. Then, invoking the approximate local optimality of the selection made by Algorithm 1 in each iteration, we can establish the desired result.

To obtain the bound in (6.19) we use the standard simple inequality (which follows from monotonicity and submodularity of  $h(\cdot)$ )

$$h(\mathcal{O}^{\text{opt}}) \leq h(\mathcal{O}) + \sum_{\underline{e} \in \mathcal{O}^{\text{opt}} \setminus \mathcal{O}} h_{\mathcal{O}}(\underline{e}) \quad (6.20)$$

Then, we can discard all elements  $\underline{e}$  from  $\mathcal{O}^{\text{opt}}$  that belong to  $\mathcal{O}^{\text{opt}} \setminus \mathcal{O}$  and for whom  $h_{\mathcal{O}}(\underline{e}) = 0$ . Let  $\tilde{\mathcal{O}}^{\text{opt}}$  denote the set after expurgating such elements from  $\mathcal{O}^{\text{opt}}$  and note that

$$h(\mathcal{O}^{\text{opt}}) \leq h(\mathcal{O}) + \sum_{\underline{e} \in \tilde{\mathcal{O}}^{\text{opt}} \setminus \mathcal{O}} h_{\mathcal{O}}(\underline{e}) \quad (6.21)$$

We next invoke Lemma 6.2 which is stated in the appendix. Let

$$\mathcal{O}_j = \cup_{\ell=1}^j \mathcal{E}_\ell, \quad j = 1, \dots, J$$

with  $\mathcal{O}_0 = \emptyset$  so that  $\mathcal{O} = \mathcal{O}_J$ . Using Lemma 6.2 in (6.21) and again invoking the monotonicity and submodularity of  $h(\cdot)$ , we obtain the inequality

$$h(\mathcal{O}^{\text{opt}}) \leq h(\mathcal{O}) + \sum_{j=1}^J \sum_{\mathcal{E} \in \tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}} \setminus \mathcal{O}} h_{\mathcal{O}_{j-1}}(\mathcal{E}) \quad (6.22)$$

Then, a key observation we can deduce is that all elements in  $\cup_{\ell=j}^J \mathcal{F}_\ell$  were eligible for selection by Algorithm 1 at the  $j^{\text{th}}$  iteration. This assertion follows from the construction of the sets  $\mathcal{F}_\ell, \ell = 1, \dots, J$ . Invoking Proposition 6.1 we know that the bit loading determined from Algorithm 2 for each possible user and mode combination has at-least  $\eta$  optimality. Consequently,

$$h_{\mathcal{O}_{j-1}}(\mathcal{E}_j) \geq \eta h_{\mathcal{O}_{j-1}}(\mathcal{E}), \quad \forall \mathcal{E} \in \cup_{\ell=j}^J \mathcal{F}_\ell \quad (6.23)$$

Using (6.32) we see that  $\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}} \setminus \mathcal{O} \subseteq \cup_{\ell=j}^J \mathcal{F}_\ell$  with  $|\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}} \setminus \mathcal{O}| \leq \lambda + 1$ . Using these observations along with (6.23) in (6.22), yields that

$$\begin{aligned} h(\mathcal{O}^{\text{opt}}) &\leq h(\mathcal{O}) + \sum_{j=1}^J \sum_{\mathcal{E} \in \tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}} \setminus \mathcal{O}} \frac{h_{\mathcal{O}_{j-1}}(\mathcal{E}_j)}{\eta} \\ &\leq h(\mathcal{O}) + \sum_{j=1}^J \frac{(\lambda + 1) h_{\mathcal{O}_{j-1}}(\mathcal{E}_j)}{\eta} = h(\mathcal{O}) \left( 1 + \frac{\lambda + 1}{\eta} \right) \end{aligned}$$

which is the desired result.  $\square$

**Theorem 6.2** *Algorithm 1 yields an  $\Omega\left(\frac{1}{\ln(\min\{N, K\}) - \ln(\ln(\min\{N, K\}))}\right)$  approximation guarantee for the problem in (6.17) (or (6.5)).*

**Proof:** We will prove this theorem by first showing that

$$h(\hat{\mathcal{O}}) \geq \frac{\gamma h(\mathcal{O})}{\mathbf{a}_{\mathcal{O}} + 1} \quad (6.24)$$

where  $\gamma \in (0, 1)$  is a constant. We next prove that  $\mathbf{a}_{\mathcal{O}}$  (which represents the sum cost after stage one) scales at-most logarithmically.

Suppose (6.24) holds. Then, define a sequence  $S_k, k = 0, 1, \dots, J$ , where  $J$  is the number of iterations in stage-1, such that  $S_0 = 0$  and  $S_k = \sum_{n=1}^N a_n \mathbf{1}(n \in \{\underline{e}_1, \dots, \underline{e}_k\})$ ,  $\forall k = 1, \dots, J$ . Thus,  $w_k = \lambda^{S_k}$ ,  $\forall k = 1, \dots, J$ . Also,  $\mathcal{O} = \{\underline{e}_1, \dots, \underline{e}_J\}$  and  $\mathbf{a}_{\mathcal{O}} = S_J = \sum_{n=1}^N a_n \mathbf{1}(n \in \{\underline{e}_1, \dots, \underline{e}_J\})$ . Notice that  $S_k - S_{k-1} \leq \mathbf{a}_{\underline{e}_k}$ ,  $k = 1, \dots, J$  and that each  $\underline{e}_k$  being feasible must satisfy  $\mathbf{a}_{\underline{e}_k} \leq 1$ ,  $\forall k \geq 1$  so that  $S_1 - S_0 = S_1 \leq 1 \leq \lambda$ . In addition, each  $\underline{e}_k$  must satisfy the (soft) constraint  $S_k - S_{k-1} \leq \frac{\lambda}{\lambda^{S_{k-1}}}$ ,  $\forall k \geq 2$ . Then, upon defining  $S_k = S_J$ ,  $\forall k > J$ , we can deduce that  $\{S_k\}_{k=0}^{\infty}$  represents a sequence in the family  $\underline{\mathcal{S}}(\lambda)$  defined in Lemma 6.3 stated in the appendix. Next, we note that in this sequence the number of strictly positive increments  $S_k - S_{k-1}$ ,  $\forall k \geq 1$  can be at-most  $\min\{N, \bar{K}\}$  (recall that  $\bar{K} \leq K$ ). Using this observation and invoking Lemma 6.3 together with Proposition 6.2, we can conclude that the theorem is true.

It remains thus to prove (6.24). To do so, we notice first that  $h(\mathcal{O}) = \sum_{n \in \mathcal{R}'} V_n$ . Further, let  $\vartheta$  be the approximation guarantee of the method employed in Step 21 (in stage-2) of Algorithm 1. For instance, simple greedy type methods for the standard knapsack problem in (6.9) yield  $\vartheta = 1/2$  [74]. Therefore,

$$h(\hat{\mathcal{O}}) \geq \vartheta \sum_{n \in \mathcal{A}} V_n, \forall \mathcal{A} \subseteq \mathcal{R}' : \sum_{n \in \mathcal{A}} a_n \leq 1. \quad (6.25)$$

Then, suppose that all the RBs in  $\mathcal{R}'$  can be divided into  $L$  bins such that the sum of cost of the RBs in each bin is less than or equal to one. We can immediately infer using (6.25) that  $h(\hat{\mathcal{O}}) \geq \frac{\vartheta h(\mathcal{O})}{L}$ . This is because the profit over each bin (i.e., sum of  $\{V_n\}$  over RBs in that bin) cannot exceed  $h(\hat{\mathcal{O}})/\vartheta$  and the sum profit over all  $L$  bins is equal to  $h(\mathcal{O}) = \sum_{n \in \mathcal{R}'} V_n$ . A simple bound on such  $L$  can be derived using the standard bin packing argument. In particular, there must be only one bin for which the sum cost is less than or equal to  $1/2$ , else we could combine two such bins into one bin. As a result  $L - 1$  bins must have their respective costs no less than  $1/2$ . Thus, the total cost of all the RBs in  $\mathcal{R}'$  (which is the sum cost over all  $L$  bins),  $\mathbf{a}_{\mathcal{O}}$ , must be at-least  $\frac{L-1}{2}$ , from which we can deduce that  $L \leq 2\mathbf{a}_{\mathcal{O}} + 1$ . This proves (6.24).  $\square$

We remark that by leveraging the proof of Lemma 6.2 we can show that

$$w_J \leq \min\{N, \bar{K}\} \lambda^2$$

, which yields a firm bound

$$\mathbf{a}_{\mathcal{O}} = \log_{\lambda}(w_J) \leq \log_{\lambda}(\min\{N, \bar{K}\}) + 2 \log_{\lambda}(\lambda).$$

Utilizing this bound in the proof of Theorem 7.2 and invoking Proposition 6.2, we can deduce that the approximation guarantee of Algorithm 1 is at-least

$$\frac{\eta^{\vartheta}}{(\eta + \lambda + 1)(2 \log_{\lambda}(\min\{N, \bar{K}\}) + 4 \log_{\lambda}(\lambda) + 1)}.$$

We also extended the algorithm incorporating multiple more general cost constraints that can be used to approximately solve the scheduling problem under network resource (bandwidth) sharing constraints [76]. This extension is shown to yield the same guarantee as the original version with one cost constraint.

### 6.2.1 Benchmarking

In this section we list alternative approaches to design algorithms for (6.5) or (6.17) that are more direct adaptations of known techniques.

**Alternative-1:** We let naive-greedy denote the direct adaptation of the classical greedy algorithm of [71] to the problem in (6.17). A smarter alternative is to apply the greedy method on the formulation in (6.18) which uses the submodular constraint instead of the knapsack one. In particular, given the elements  $\underline{e}_1^g, \dots, \underline{e}_{j-1}^g$  selected so far, where  $\underline{e}_0^g = \emptyset$ , we simply choose the next one,  $\underline{e}_j^g$ , as

$$\begin{aligned} \arg \max_{\underline{e} \in \Psi \setminus \{\underline{e}_1^g, \dots, \underline{e}_{j-1}^g\}} \{ & h(\{\underline{e}_1^g, \dots, \underline{e}_{j-1}^g, \underline{e}\}) - h(\{\underline{e}_1^g, \dots, \underline{e}_{j-1}^g\}) \} \\ \text{s.t. } & \{\underline{e}_1^g, \dots, \underline{e}_{j-1}^g, \underline{e}\} \in \mathcal{I}, \\ & \sum_{n \in \mathcal{N}} a_n \mathbf{1} \left( n \in \{\underline{e}_1^g, \dots, \underline{e}_{j-1}^g, \underline{e}\} \right) \leq 1. \end{aligned} \quad (6.26)$$



This problem can in turn be approximately solved using Algorithm 2. The process terminates if no such  $\underline{e}_j^g$  can be found. Invoking the fact that  $h(\cdot)$  is also sub-additive and that  $\underline{e}_1^g$  is the (approximately) optimal choice among all (singleton) elements in  $\Psi$ , i.e.,  $h(\underline{e}_1^g) \geq \eta h(\underline{e})$ ,  $\forall \underline{e} \in \Psi$ , we can show that Alternative-1 yields an approximation factor  $\Omega\left(\frac{1}{\min\{N, K\}}\right)$ .

**Alternative-2:** This alternative first models the matroid constraint in (6.17) as  $K + 1$  knapsack constraints. Indeed, out of these  $K + 1$  constraints, the first  $K$  constraints enforce the fact that each user can be selected at-most once whereas the last one enforces the user limit (cardinality) constraint. Then, accounting for the sum cost constraint we have a submodular maximization problem subject to  $K + 2$  linear constraints. Direct adaptation of the multiplicative updates algorithm of [72] yields an approximation ratio  $\Omega(1/K)$ . We note here that although the first  $K + 1$  constraints are binary and sparse, the last constraint is not binary valued so that the column sparse multiplicative updates algorithm from [72] is not applicable.

Comparing these two alternatives with Algorithm 1, we see that the latter one achieves a significantly better approximation ratio. Finally, we provide upper bounds to bound the optimality gap of Algorithm 1 for every input instance. For the backlogged traffic model, we note that the constraints of (6.4) involving the max function can be easily converted into linear inequality constraints so the problem in (6.4) is indeed a integer linear programming problem (ILP). Upon relaxing (6.4) by allowing  $\{x_{u,n}^m \in [0, 1]\}$  we can obtain the LP upper bound. Considering the general finite queue model, we note that (6.5) is not an ILP. Nevertheless, specializing to the 0 – 1 throughput

model we can obtain the following LP upper bound.

$$\begin{aligned}
& \max_{\{x_{u,n}^m\}} \sum_{u=1}^K \sum_{n=1}^N \sum_{m=1}^M \alpha_u x_{u,n}^m B_{u,n}^m \\
& \text{subject to } \sum_{m=1}^M \sum_{u=1}^K x_{u,n}^m \leq 1, \quad 1 \leq n \leq N \\
& \sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq 1, \quad 1 \leq u \leq K \\
& \sum_{u=1}^K \sum_{m=1}^M \max_{1 \leq n \leq N} x_{u,n}^m \leq \bar{K} \\
& \sum_{n=1}^N \sum_{u=1}^K \sum_{m=1}^M a_n x_{u,n}^m \leq 1 \\
& \sum_{m=1}^M \sum_{n=1}^N x_{u,n}^m B_{u,n}^m \leq Q_u, \quad 1 \leq u \leq K \\
& x_{u,n}^m \in [0, 1] \quad \forall u, m, n,
\end{aligned} \tag{6.27}$$

### 6.3 Simulation Results

We conducted a systematic evaluation of our proposed algorithms over a homogeneous single-cell system and the 0 – 1 throughput model. To outer bound the performance we employed the LP upper bound (6.27), whereas for comparison with other schemes yielding an achievable weighted sum throughput we used the naive greedy as well as the Alternative-1 algorithms. For practical considerations, we restricted our attention to these latter two schemes since they are both deterministic and have a low complexity comparable to Algorithm 1. In our study, we chose the number of modes to be  $M = 4$ . We also fixed  $\lambda = \theta = 10$  when implementing Algorithms 1 and 2, respectively, and implemented the pruning step of Algorithm 1 in a greedy manner. Further, we considered  $a_n = \frac{1}{\lceil N/\ln(N) \rceil}$ ,  $n \in \mathcal{N}$  so that the cost constraint mandates that no more than  $\ln(N)$  RBs out of the  $N$  can be allocated. Under the 0 – 1 model, the maximal bit loadings,  $\{B_{u,n}^m\}$ , as well as the queue sizes were generated using the half-normal distribution. We remark that the latter two choices were observed to capture *harder* input instances where the gap to the upper bound for our algorithm was relatively

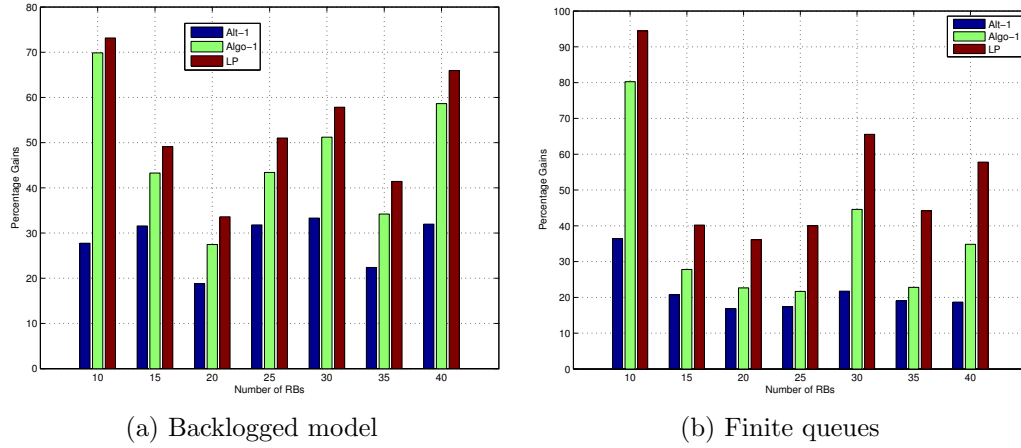
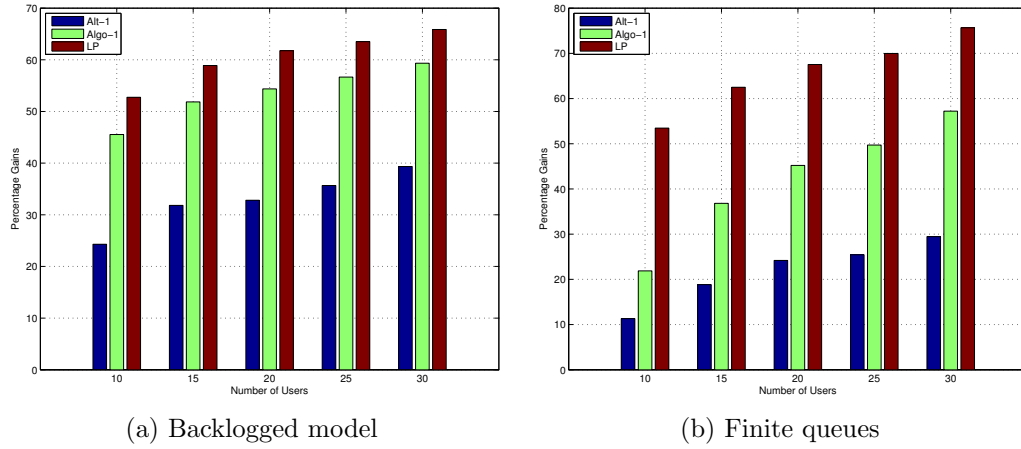
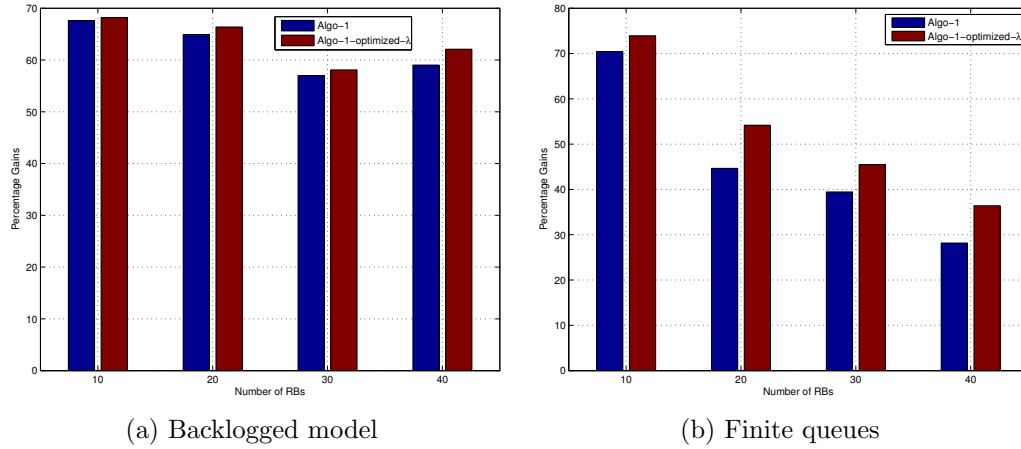
larger. Other choices for generating the maximal bit loadings (such as using Rayleigh fading and the Shannon formula) were seen to result in smaller gaps.

We obtained average performance over 500 input traces considering following two scenarios:

**Scenario I:** The number of RBs,  $N$ , was varied between 10 to 40, while keeping the number of users and the user limit fixed at,  $K = \bar{K} = 20$ . The user weights were chosen to be identical to unity. The percentage gains obtained by Alternative-1 and our Algorithm 1 as well the LP bound over the naive greedy method are depicted in Figures 6.1a and 6.1b, which cover the cases with large and small queue size regimes, respectively.

**Scenario II:** The number of users,  $K$ , was varied between 10 to 30, with the user limit  $\bar{K}$  set to  $\lfloor \frac{2K}{3} \rfloor$ , while keeping the number of RBs fixed at,  $N = 40$ . User weights were generated randomly using the uniform distribution, independently across traces. The percentage gains obtained over the naive greedy method are depicted in Figures 6.2a and 6.2b, which again cover the cases with large and small queue size regimes, respectively.

From the figures, we see that substantial gains can be achieved by our proposed algorithm. The gap to optimal seems larger in the regime with the small queue sizes. One reason could be that Algorithm 2 can introduce a larger sub-optimality in solving the sub-problems since the queue constraints are active more often. However, we caution that the LP bound itself can be loose in this regime. Finally, we evaluated the performance of Algorithm 1 for different values of  $\lambda$  in Figs. 6.3a and 6.3b ( $\theta$  was fixed at 10 as before) over the first scenario described above. For each trace we implemented Algorithm 1 for 10 different values of  $\lambda$  and selected the best one to obtain the optimized result. The other result always uses the choice  $\lambda = 10$  as before. We see that good gains are obtained in the small queue sizes regime. More importantly, the choice  $\lambda = 2$  was observed to be the best choice for fixed  $\lambda$  and indeed fixing  $\lambda = 2$  resulted in almost the same performance as that shown in Figs. 6.3a and 6.3b when the best  $\lambda$  was chosen for each trace. A systematic way to optimize and determine good (fixed) choices of  $\lambda$  and  $\theta$  is an open problem.

Figure 6.1: Weighted sum rate vs number of RBs for  $\lambda = 10$ ,  $\theta = 10$ Figure 6.2: Weighted sum rate vs number of users for  $\lambda = 10$ ,  $\theta = 10$ Figure 6.3: Comparison of weighted sum rate vs number of RBs for the case with  $\lambda = 10$  and  $\theta = 10$ , and the case with the optimized  $\lambda$  and  $\theta = 10$

## 6.4 Conclusions

We proposed a novel algorithm for energy-aware subframe scheduling over the LTE downlink. The proposed algorithm is simple enough to be implementable and we rigorously established its performance guarantee. The main advantage of the algorithm is that it can lead to substantial improvements over other greedy algorithms of comparable complexity, both in terms of a provable approximation guarantee as well as average case performance obtained in simulations. We note that in certain special cases when cost constraints involve identical prices, the constraint structure can be exploited to in-fact achieve constant-factor approximations.

## Appendix

**Definition 6.1** Let  $\Psi$  be a ground set and  $h : 2^\Psi \rightarrow \mathbb{R}_+$  be a non-negative set function defined on the subsets of  $\Psi$ , that is also normalized ( $h(\emptyset) = 0$ ) and non-decreasing ( $h(\mathcal{A}) \leq h(\mathcal{B})$ ,  $\forall \mathcal{A} \subseteq \mathcal{B} \subseteq \Psi$ ). Then, the set function  $h(\cdot)$  is a submodular set function if it satisfies,

$$h(\mathcal{B} \cup a) - h(\mathcal{B}) \leq h(\mathcal{A} \cup a) - h(\mathcal{A}), \quad \forall \mathcal{A} \subseteq \mathcal{B} \subseteq \Psi \text{ \& } a \in \Psi \setminus \mathcal{B}.$$

Next, for any given subset  $\mathcal{B} \subseteq \Psi$ ,  $\{\eta_q, \mathcal{T}_q\}$  is said to be a fractional cover of  $\mathcal{B}$ , if  $\eta_q \in [0, 1]$ ,  $\mathcal{T}_q \subseteq \Psi \quad \forall q$  and  $\sum_{q: a \in \mathcal{T}_q} \eta_q \geq 1$ ,  $\forall a \in \mathcal{B}$ . The set function  $h(\cdot)$  is fractionally sub-additive if for any given subset  $\mathcal{B} \subseteq \Psi$  and its fractional cover  $\{\eta_q, \mathcal{T}_q\}$ , it satisfies

$$h(\mathcal{B}) \leq \sum_q \eta_q h(\mathcal{T}_q).$$

Note that considering all such normalized and non-decreasing set functions, the class of fractionally sub-additive set function subsumes that of submodular set functions and is subsumed in turn by the class of sub-additive set functions [77].

**Definition 6.2**  $(\Omega, \underline{I})$ , where  $\underline{I}$  is collection of some subsets of  $\Omega$ , is said to be a matroid if  $\underline{I}$  is an independence family, i.e.,

- $\underline{I}$  is downward closed, i.e.,  $\mathcal{A} \in \underline{I} \text{ \& } \mathcal{B} \subseteq \mathcal{A} \Rightarrow \mathcal{B} \in \underline{I}$
- For any two members  $\mathcal{F}_1 \in \underline{I}$  and  $\mathcal{F}_2 \in \underline{I}$  such that  $|\mathcal{F}_1| < |\mathcal{F}_2|$ , there exists  $e \in \mathcal{F}_2 \setminus \mathcal{F}_1$  such that  $\mathcal{F}_1 \cup \{e\} \in \underline{I}$ . This property is referred to as the exchange property.

**Proof:** ( **Theorem 6.1**) We first show that for the 0 – 1 throughput model (6.3) the problem in (6.5) can be cast in the form of a combinatorial auction problem (a.k.a. welfare maximization problem) where items (RBs) in  $\mathcal{N}$  have to be assigned in a non-overlapping manner to the  $K$  bidders (users) and each valuation function is given by a set function,  $h_u : 2^{\mathcal{N}} \rightarrow \mathbb{R}_+$ . To do so, we define

$$h_u(\mathcal{R}) = \alpha_u \max_m \{ \min \{ Q_u, \sum_{n \in \mathcal{R}} B_{u,n}^m \} \}, \mathcal{R} \subseteq \mathcal{N}.$$

Then, we can express (6.5) as

$$\begin{aligned} & \max_{\{\mathcal{S}_u \subseteq \mathcal{N}\}_{u=1}^K} \sum_{u=1}^K \{h_u(\mathcal{S}_u)\} \\ & \text{s.t.} \sum_{u=1}^K \mathbf{1}(\ell \in \mathcal{S}_u) \leq 1, \forall \ell \in \mathcal{N}, \\ & \sum_{u=1}^K \sum_{\ell \in \mathcal{S}_u} a_\ell \leq 1. \end{aligned} \tag{6.28}$$

Note that in (6.28) unlike the standard form there is an additional linear cost constraint on the items in (6.5). Note that for each user  $u$ , the set function  $h_u(\cdot)$  need not be submodular even in the backlogged model. Invoking the definition given above, for any set  $\mathcal{S} \subseteq \mathcal{N}$  and any fractional cover  $\{\eta_q, \mathcal{T}_q\}$  of  $\mathcal{S}$ , we have to prove  $h_u(\mathcal{S}) \leq \sum_q \eta_q h_u(\mathcal{T}_q)$ . We proceed by assuming, without loss of generality, that the user weights are unity  $\alpha_u = 1, \forall u$ . Then, let  $m^*$  be the mode that is optimal for the user  $u$  and set  $\mathcal{S}$ , i.e.

$$h_u(\mathcal{S}) = \min \left\{ Q_u, \sum_{n \in \mathcal{S}} B_{u,n}^{m^*} \right\}. \tag{6.29}$$

Suppose that  $h_u(\mathcal{S}) = \sum_{n \in \mathcal{S}} B_{u,n}^{m^*} \leq Q_u$ . Then,

$$h_u(\mathcal{T}_q) \geq \min \left\{ Q_u, \sum_{n \in \mathcal{T}_q} B_{u,n}^{m^*} \right\} \geq \min \left\{ Q_u, \sum_{n \in \mathcal{T}_q \cap \mathcal{S}} B_{u,n}^{m^*} \right\} = \sum_{n \in \mathcal{T}_q \cap \mathcal{S}} B_{u,n}^{m^*}. \tag{6.30}$$

It follows that

$$\begin{aligned}
\sum_q \eta_q h_u(\mathcal{T}_q) &\geq \sum_q \eta_q \sum_{n \in \mathcal{T}_q \cap \mathcal{S}} B_{u,n}^{m*} \\
&= \sum_{n \in \mathcal{S}} B_{u,n}^{m*} \underbrace{\sum_{q: n \in \mathcal{T}_q} \eta_q}_{\geq 1} \\
&\geq \sum_{n \in \mathcal{S}} B_{u,n}^{m*} = h_u(\mathcal{S}),
\end{aligned}$$

which settles the result for this case. It remains to prove the result when  $h_u(\mathcal{S}) = Q_u \leq \sum_{n \in \mathcal{S}} B_{u,n}^{m*}$ . In this case we can find a subset  $\mathcal{R} \subseteq \mathcal{S}$  such that  $\sum_{n \in \mathcal{R}} B_{u,n}^{m*} \geq Q_u$  but all its strict subsets  $\mathcal{A} \subset \mathcal{R}$  satisfy  $\sum_{n \in \mathcal{A}} B_{u,n}^{m*} < Q_u$ . Upon obtaining such an  $\mathcal{R}$ , we can partition the cover  $\{\mathcal{T}_q\}$  into two parts  $\{\mathcal{T}_q\}_{q \in \mathcal{I}_1}$  and the remaining sets of the cover are in  $\{\mathcal{T}_q\}_{q \in \mathcal{I}_2}$ . Since  $\sum_{n \in \mathcal{T}_q} B_{u,n}^{m*} \geq Q_u \forall q \in \mathcal{I}_1$ , we have  $h_u(\mathcal{T}_q) = Q_u \forall q \in \mathcal{I}_1$ . Consequently,

$$\begin{aligned}
\sum_q \eta_q h_u(\mathcal{T}_q) &\geq Q_u \underbrace{\sum_{q: q \in \mathcal{I}_1} \eta_q}_{\triangleq \beta} + \sum_{q: q \in \mathcal{I}_2} \eta_q \sum_{n \in \mathcal{T}_q \cap \mathcal{R}} B_{u,n}^{m*} \\
&= Q_u \beta + \sum_{n \in \mathcal{R}} B_{u,n}^{m*} \sum_{q \in \mathcal{I}_2: n \in \mathcal{T}_q} \eta_q \tag{6.31}
\end{aligned}$$

Notice that if  $\beta \geq 1$  (which must always hold when  $\mathcal{R}$  is a singleton set) the desired inequality is already proved. On the other hand, if  $\beta < 1$  then since  $\{\eta_q, \mathcal{T}_q\}$  is a fractional cover of  $\mathcal{S}$ , we can deduce that for each  $n \in \mathcal{R}$ ,

$$\sum_{q \in \mathcal{I}_2: n \in \mathcal{T}_q} \eta_q \geq 1 - \sum_{q \in \mathcal{I}_1: n \in \mathcal{T}_q} \eta_q \geq 1 - \beta,$$

which using (6.31) and the fact that  $\sum_{n \in \mathcal{R}} B_{u,n}^{m*} \geq Q_u$ , yields the desired result.  $\square$

**Lemma 6.2** *The optimal solution (after expurgation),  $\tilde{\mathcal{O}}^{\text{opt}}$ , can be partitioned as  $\tilde{\mathcal{O}}^{\text{opt}} = \tilde{\mathcal{O}}_{\mathcal{I}_1}^{\text{opt}} \cup \dots \cup \tilde{\mathcal{O}}_{\mathcal{I}_J}^{\text{opt}}$ , where*

$$\begin{aligned}
\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}} &\subseteq \cup_{\ell=j}^J \tilde{\mathcal{F}}_{\ell}, \\
|\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}}| &\leq \lambda + 1, \forall j = 1, \dots, J \tag{6.32}
\end{aligned}$$

**Proof:** First note that there can be at-most  $\bar{K} \leq K$  elements in the set  $\tilde{\mathcal{O}}^{\text{opt}}$  and all those elements must have distinct users, since  $\tilde{\mathcal{O}}^{\text{opt}}$  has to satisfy the matroid constraint



in (6.17). Further, without loss of optimality, we can assume that  $\tilde{\mathcal{O}}^{\text{opt}}$  also satisfies the (stricter) linear knapsack constraint in (6.17) so that  $\sum_{\underline{e} \in \tilde{\mathcal{O}}^{\text{opt}}} \mathbf{a}_{\underline{e}} \leq 1$ . Then, to prove this lemma it suffices to show that

$$|\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})| \leq j(\lambda + 1), \forall j = 1, \dots, J. \quad (6.33)$$

Indeed, suppose that (6.33) holds. Then, one way of constructing the sets  $\{\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}}\}_{j=1}^J$  satisfying (6.32) is as follows. All members of  $\tilde{\mathcal{O}}^{\text{opt}} \cap \tilde{\mathcal{F}}_1$  are included in  $\tilde{\mathcal{O}}_{\mathcal{I}_1}^{\text{opt}}$ , which we note from (6.33) are at-most  $\lambda + 1$  in number. Then, all elements from  $\tilde{\mathcal{O}}^{\text{opt}} \cap \tilde{\mathcal{F}}_2$  are included in  $\tilde{\mathcal{O}}_{\mathcal{I}_1}^{\text{opt}}$ , followed by all the elements from  $\tilde{\mathcal{O}}^{\text{opt}} \cap \tilde{\mathcal{F}}_3$  and so on, until the cardinality of  $\tilde{\mathcal{O}}_{\mathcal{I}_1}^{\text{opt}}$  reaches  $\lambda + 1$ . Note that if only a partial subset of elements from  $\tilde{\mathcal{O}}^{\text{opt}} \cap \tilde{\mathcal{F}}_{\ell}$  (for some  $\ell \geq 2$ ) can be included without violating the cardinality bound  $\lambda + 1$ , then any arbitrary subset of  $\tilde{\mathcal{O}}^{\text{opt}} \cap \tilde{\mathcal{F}}_{\ell}$  which ensures  $|\tilde{\mathcal{O}}_{\mathcal{I}_1}^{\text{opt}}| = \lambda + 1$  can be included. The construction of  $\tilde{\mathcal{O}}_{\mathcal{I}_j}^{\text{opt}}$   $j \geq 2$ , begins once  $\lambda + 1$  elements have been included in  $\tilde{\mathcal{O}}_{\mathcal{I}_{j-1}}^{\text{opt}}$ , and proceeds in a similar manner considering all the remaining elements  $\tilde{\mathcal{O}}^{\text{opt}} \setminus (\cup_{\ell=1}^{j-1} \tilde{\mathcal{O}}_{\mathcal{I}_{\ell}}^{\text{opt}})$ . If at any  $j$  all elements in  $\tilde{\mathcal{O}}^{\text{opt}}$  have been included, then the sets  $\{\tilde{\mathcal{O}}_{\mathcal{I}_{\ell}}^{\text{opt}}\}_{\ell=j+1}^J$  are simply defined to be empty sets. The key point is that at any step  $j \geq 2$ , due to (6.33) the set of remaining elements  $\tilde{\mathcal{O}}^{\text{opt}} \setminus (\cup_{\ell=1}^{j-1} \tilde{\mathcal{O}}_{\mathcal{I}_{\ell}}^{\text{opt}})$ , will not have any elements from  $\cup_{\ell=1}^{j-1} \tilde{\mathcal{F}}_{\ell}$  which ensures that the construction satisfies (6.32).

We next prove (6.33) via contradiction. Suppose first that  $\exists j \in \{1, \dots, J-1\}$  for which

$$|\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})| > j(\lambda + 1). \quad (6.34)$$

Note that by construction  $\mathcal{O} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell}) = \{\underline{e}_1, \dots, \underline{e}_j\}$  so that  $|\mathcal{O} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})| = j$ . Considering the set  $\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})$ , collect all elements which have an identical user as some element in  $\{\underline{e}_1, \dots, \underline{e}_j\}$ . There can be at-most  $j$  such elements since each element in  $\tilde{\mathcal{O}}^{\text{opt}}$  must have a distinct user. Let  $\mathcal{M}_1^j$  denote the set obtained by expurgating all such elements from  $\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})$ . Thus, we have that

$$|\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_{\ell})| \leq j + |\mathcal{M}_1^j|. \quad (6.35)$$

Next, consider the set  $\mathcal{M}_1^j$ . Each element in this set must belong to one of the sets  $\mathcal{F}_\ell$ ,  $\ell = 1, \dots, j$ . Indeed, each  $\underline{e} \in \mathcal{M}_1^j$  was assigned to some  $\mathcal{F}_\ell$ ,  $\ell \in \{1, \dots, j\}$  because, it satisfied  $(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell, \underline{e}\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell\}})w_\ell > \lambda$ . We remind that such an element could not have been declared infeasible on account of sharing a common user (due to construction of  $\mathcal{M}_1^j$ ) or on account of user limit being reached (since  $j \leq J - 1$ ). Then, notice that  $\mathbf{a}_{\underline{e}} \geq \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell, \underline{e}\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell\}}$ ,  $\forall \ell$  and that the sequence  $\{w_\ell\}_{\ell=1}^j$  is non-decreasing, i.e.,  $w_\ell \leq w_j, \forall \ell \in \{1, \dots, j\}$ . Therefore,

$$|\mathcal{M}_1^j| \leq \sum_{\underline{e} \in \mathcal{M}_1^j} \frac{\mathbf{a}_{\underline{e}} w_j}{\lambda}. \quad (6.36)$$

Further, since  $\mathcal{M}_1^j \subseteq \tilde{\mathcal{O}}^{\text{opt}}$  we must have  $\sum_{\underline{e} \in \mathcal{M}_1^j} \mathbf{a}_{\underline{e}} \leq 1$ . Using this fact in (6.36) and substituting in (6.35), we obtain

$$|\tilde{\mathcal{O}}^{\text{opt}} \cap (\cup_{\ell=1}^j \tilde{\mathcal{F}}_\ell)| \leq j + \frac{w_j}{\lambda}. \quad (6.37)$$

Finally, we can expand  $\frac{w_j}{\lambda}$  as

$$\frac{w_j}{\lambda} = \frac{w_1}{\lambda} + \underbrace{\frac{w_2}{\lambda} - \frac{w_1}{\lambda}} + \dots, \underbrace{\frac{w_j}{\lambda} - \frac{w_{j-1}}{\lambda}}. \quad (6.38)$$

Considering each term in the RHS of (6.38) we see that  $\frac{w_1}{\lambda} = \frac{\lambda^{\mathbf{a}_{\underline{e}_1}}}{\lambda} \leq 1 \leq \lambda$ , since  $\mathbf{a}_{\underline{e}_1} \leq 1$  and  $\lambda > 1$ . On the other hand, for  $2 \leq \ell \leq j$ ,

$$\begin{aligned} \frac{w_\ell}{\lambda} - \frac{w_{\ell-1}}{\lambda} &= \frac{w_{\ell-1}(\lambda^{\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_{\ell-1}\}}} - 1)}{\lambda} \\ &\leq \frac{w_{\ell-1}(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_{\ell-1}\}})\lambda}{\lambda} \\ &\leq \frac{\lambda^2}{\lambda} = \lambda. \end{aligned} \quad (6.39)$$

The penultimate inequality follows from the fact that  $y^x - 1 \leq xy$ ,  $\forall y \geq 0$  and  $x \in [0, 1]$ . The last inequality follows from that fact that  $\underline{e}_\ell$  (for any  $\ell$ ) was added after verifying that it satisfied the soft constraint

$$w_{\ell-1}(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_\ell\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_{\ell-1}\}}) \leq \lambda.$$

Since there are exactly  $j$  terms in the RHS of (6.38) and as shown in (6.39) each term is upper-bounded by  $\lambda$ , we see that  $\frac{w_j}{\lambda} \leq j\lambda$ . We can now use (6.37) to obtain the desired

contradiction of (6.34) and conclude that (6.33) is true for all  $j = 1, \dots, J-1$ . Lastly, at  $j = J$ , if  $J = \bar{K}$  then the desired result in (6.33) is trivially true. If  $J < \bar{K}$  then the stage-1 of the algorithm did not terminate on account of user limit being reached. Thus, any element  $\underline{e} \in \tilde{\mathcal{O}}^{\text{opt}}$  that lies in  $\mathcal{F}_J$  and does not share a common user with  $\underline{e}_J$  must satisfy  $(\mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_J, \underline{e}\}} - \mathbf{a}_{\{\underline{e}_1, \dots, \underline{e}_J\}})w_J > \lambda$ , since such an element cannot have an incremental gain  $h_{\mathcal{O}}(\underline{e}) = 0$  (due to construction of  $\tilde{\mathcal{O}}^{\text{opt}}$ ). The argument provided above can now be directly reused to conclude that (6.33) is true for  $j = J$  as well.  $\square$

**Lemma 6.3** *Let  $\underline{\mathcal{S}}(\lambda)$  denote a family of non-negative and non-decreasing sequences parameterized by the scalar  $\lambda > 1$ . Each sequence  $\{S_k\}_{k=0}^\infty$  in  $\underline{\mathcal{S}}(\lambda)$  must satisfy  $S_0 = 0$  and  $S_{k+1} - S_k \leq \frac{\lambda}{\lambda^{S_k}}, \forall k \geq 0$ . Then, given any positive scalar  $\Delta$  such that  $\Delta \ln(\lambda) > 1$ , there exists a finite integer  $\hat{k} \geq 1$ , such that  $S_k \leq \Delta(\ln(k) - \ln(\ln(k))), \forall k \geq \hat{k}$  holds true for all sequences in the family  $\underline{\mathcal{S}}(\lambda)$ .*

**Proof:** Let  $\Delta$  be any positive scalar such that  $\alpha = \Delta \ln(\lambda) > 1$ . Suppose there exists a sequence  $\{S_k\}_{k=0}^\infty$  in  $\underline{\mathcal{S}}(\lambda)$  which satisfies  $S_k \geq \Delta(\ln(k) - \ln(\ln(k))), \forall k \geq k'$  for some  $k' \geq 1$ . Such a sequence is clearly diverging. However, since it belongs to the family  $\underline{\mathcal{S}}(\lambda)$ , we must have that the increments  $\delta_{k+1} = S_{k+1} - S_k, \forall k$  must satisfy  $\delta_{k+1} \leq \frac{\lambda}{\lambda^{S_k}} \leq \frac{\lambda(\ln(k))^\alpha}{k^\alpha}, \forall k \geq k'$ . Then, since  $\alpha > 1$  we must have that the sum of increments converges, i.e.,  $\sum_{k:k \geq k'} \delta_k < \infty$  which is a contradiction. Therefore for any series  $\{S_k\}_{k=0}^\infty$  in  $\underline{\mathcal{S}}(\lambda)$ , we cannot have a positive integer  $k'$  such that  $S_k \geq \Delta(\ln(k) - \ln(\ln(k))), \forall k \geq k'$ . We can obtain a sharper observation that since  $S_{k'}$  is finite (indeed  $S_k \leq k\lambda, \forall k$ ) and  $S_{k'+N} = S_{k'} + \sum_{j=1}^N \delta_{k'+j}$ , we cannot have  $S_k \geq \Delta(\ln(k) - \ln(\ln(k))), \forall k : k' \leq k \leq k' + N$  for all large enough  $N$ . To summarize, we can conclude that for any given integer  $k' \geq 1$ , there exists a large enough integer  $N' \geq 1$  such that no sequence in  $\underline{\mathcal{S}}(\lambda)$  can have a segment (spanning  $k', k' + 1, \dots, k' + N'$ ) which is point-wise greater than  $\Delta(\ln(k) - \ln(\ln(k))), \forall k : k' \leq k \leq k' + N$ .

Next, we show that there exists a fixed  $\check{k} \geq 1$  for which the following holds true for all sequences in  $\underline{\mathcal{S}}(\lambda)$ ,

$$\begin{aligned} S_{k+1} &\geq \Delta(\ln(k+1) - \ln(\ln(k+1))) \Rightarrow \\ S_k &\geq \Delta(\ln(k) - \ln(\ln(k))), \forall k \geq \check{k}. \end{aligned} \tag{6.40}$$

We establish (6.40) via contradiction. Suppose there exist a sequence whose pair  $S_{k+1}, S_k$  is such that  $S_{k+1} = \Delta(\ln(k+1) - \ln(\ln(k+1))) + \theta$  &  $S_k = \Delta(\ln(k) - \ln(\ln(k))) - \eta$  for some  $\eta \geq 0$  and  $\theta \geq 0$ . Using the fact that  $\delta_{k+1} \leq \frac{\lambda}{\lambda^{S_k}} = \frac{\lambda^{\eta+1}(\ln(k))^\alpha}{k^\alpha}$  we have that

$$\begin{aligned} \delta_{k+1} &= \Delta \ln \left( \frac{(k+1) \ln(k)}{k \ln(k+1)} \right) + \theta + \eta \leq \frac{\lambda^{\eta+1}(\ln(k))^\alpha}{k^\alpha} \Rightarrow \\ &\left( \frac{k}{\ln(k)} \right)^\alpha (\Delta \ln \left( \frac{(k+1) \ln(k)}{k \ln(k+1)} \right) + \theta + \eta) \leq \lambda^{\eta+1} \end{aligned} \quad (6.41)$$

Then since all increments are bounded above by  $\lambda$ , we must have that  $\eta \leq \lambda$ . Thus a necessary condition for (6.41) to be true is

$$\left( \frac{k}{\ln(k)} \right)^\alpha \Delta \ln \left( \frac{(k+1) \ln(k)}{k \ln(k+1)} \right) \leq \lambda^{\lambda+1} \quad (6.42)$$

However, it can be verified that since  $\alpha > 1$ ,  $\lim_{k \rightarrow \infty} \left( \frac{k}{\ln(k)} \right)^\alpha \Delta \ln \left( \frac{(k+1) \ln(k)}{k \ln(k+1)} \right) = \infty$ . Thus, there must exist a sufficiently large  $\check{k}$  such that  $\left( \frac{k}{\ln(k)} \right)^\alpha \Delta \ln \left( \frac{(k+1) \ln(k)}{k \ln(k+1)} \right) > \lambda^{\lambda+1}$ ,  $\forall k \geq \check{k}$ . This yields the desired contradiction and establishes (6.40). The implication from (6.40) is that if any sequence in  $\underline{\mathcal{S}}(\lambda)$  satisfies  $S_{\check{k}+N} \geq \Delta(\ln(\check{k}+N) - \ln(\ln(\check{k}+N)))$  for any  $N \geq 1$ , then that sequence must also satisfy  $S_k \geq \Delta(\ln(k) - \ln(\ln(k)))$ ,  $\forall k : \check{k} \leq k \leq \check{k} + N$ .

Finally, we can combine these above two results to prove the lemma.  $\square$

## Chapter 7

### Optimizing Energy Efficiency over Energy-Harvesting LTE Cellular Networks

Deployment of energy harvesting devices to harness renewable energy sources in LTE networks offers two-fold advantages. Foremost, it will reduce the carbon footprint of these networks which are seeing an exponential growth. Secondly, it can also extend the lifetime of such networks. This is particularly vital in scenarios where easy accessibility for maintenance cannot be ensured. However, a renewable energy source providing gradual and unreliable energy supply adds an energy causality constraint to the system, which needs to be considered in the radio resource management (RRM). This subject has been studied under different scenarios and assumptions; see, cf. [3, 78] for the point-to-point channel, [7, 9] which consider networks with a small number of nodes, [2] for off-line scheduling algorithms and [5] for on-line ones. Recent works like [79, 80] focus on the joint allocation of energy and subchannels over a system where the harvested energy levels and subchannel gains can be predicted in advance for a scheduling period. Further advances have been made in [81] where online weighted sum rate maximization is considered and in [82] where optimization of the proportional fairness utility over such a system is setup as a biconvex optimization problem. A general framework for utility optimization over such networks has been recently proposed [83].

Energy efficiency has become an increasingly popular RRM paradigm in wireless networks, cf. [62] for a comprehensive overview. A popular energy efficiency metric is the system or global energy efficiency (GEE), which is defined as the ratio of the achieved weighted sum throughput and the energy consumed. Another important variant that has received wide attention is the weighted sum of individual energy efficiencies (WSEE). These two metrics have been optimized recently for the single-cell

downlink [67], followed by the multi-cell one [68]. Fractional programming has emerged as a popular tool to solve the resulting problems, which typically are formulated as continuous optimization problems.

Our goal in this chapter is to optimize energy efficiency metrics over multi-channel wireless networks that conform to the LTE standard while satisfying energy causality constraints imposed by the energy harvesting process. Optimizing energy efficiency over such networks requires us to carefully account for certain mandatory constraints placed by the LTE standard on the choice of transmission parameters, which are all discrete valued. We demonstrate that formulations which aim to optimize the GEE and the WSEE over a single-cell energy harvesting downlink, respectively, subject to all the main LTE constraints are indeed tractable approximately and thereby circumventing the need for continuous relaxation of the underlying discrete variables. We prove this by showing that a key sub-problem that aims to maximize the weighted sum rate under linear causality constraints on the set of used resource blocks (RBs) (or sub-channels) can be cast as a constrained submodular maximization problem. We remark here that submodular maximization [71] can be regarded as the analogue of concave maximization, over discrete (combinatorial) problems. Our result enables us to derive constant-factor approximation algorithms for optimizing both GEE and WSEE over LTE networks powered by energy harvesting devices, which to the best of our knowledge, are the first such algorithms. The proposed algorithms are deterministic and simple, thus amenable to implementation in base-stations. Simulations show that they readily outperform other heuristics of comparable complexities.

## 7.1 Problem formulation

We focus on a single-cell downlink in an LTE network that comprises of a base station serving  $K$  active users. Over each subframe (of duration 1 millisecond) the available bandwidth is partitioned into  $N$  time-frequency units referred to as RBs and each RB is the minimum assignable time-frequency resource unit. The base station is equipped with a battery of infinite storage capacity and a renewable energy source that provides energy  $E_\ell$  (Joules) in each subframe  $\ell$  [80]. We assume that  $E_\ell$  can be accurately

estimated in advance, which is true for a renewable source such as solar, which can be accurately predicted for durations up to one hour. We then incorporate a setup (as in [80]) where the channel estimates are known in advance for each block of  $L$  contiguous subframes, i.e., at the beginning of each block, channel estimates for all users for each one of the subframes in that block are available at the base station scheduler. While such a non-causal assumption might seem implausible, we note that in a practical system channel estimates are obtained from each user with a certain periodicity (such as 5 or 10 ms in LTE networks). A very typical assumption (especially for low user-mobility scenarios) is to treat the user channels (on each RB) as block fading with a coherence time equal to the configured periodicity. This assumption clearly fits our framework, and indeed the latter also covers interpolation schemes where the estimates for all subframes in a block are obtained at the beginning of that block, based on all the past available channel estimates as well as ACK/NACK feedback. Thus, we consider a block of  $L$  contiguous subframes and suppress dependence on the block index. We define

$$\mathcal{N} = \mathcal{N}_1 \cup \mathcal{N}_2 \cdots \mathcal{N}_L,$$

where

$$\mathcal{N}_\ell = \{(\ell - 1)N + 1, \dots, \ell N\},$$

denotes the set containing RBs in the  $\ell^{th}$ ,  $\ell = 1, \dots, L$  subframe of a block.

### 7.1.1 Practical Constraints

We consider two figures of merit pertaining to energy efficiency and optimize them over each block, under a set of important practical constraints. Constraints C1, C2, and C3 are the same as those introduced in Chapter 6 except that the constraint on the number of scheduled users  $\bar{K}$  is over the users that are assigned at least one RB over the scheduling block (instead of a subframe). Also, all the RBs assigned to the same user are assigned the same mode over a scheduling block. Due to the harvesting energy, the transmit energy is constrained differently as follows.

**[C4]:** The transmit energy,  $\rho$ , expended over each RB used for data transmission

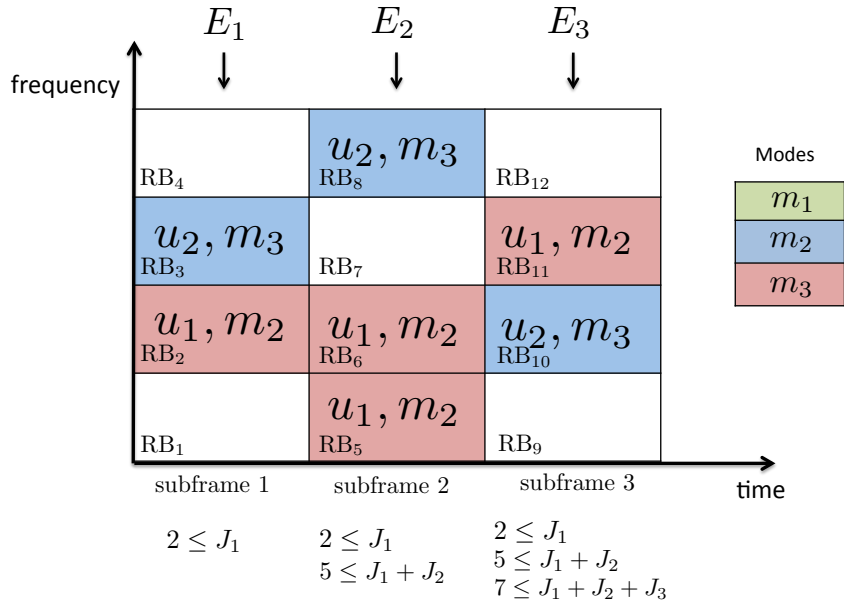


Figure 7.1: A feasible allocation for a system with  $N = 4$  RBs,  $L = 3$  subframes per block,  $K = 2$  users and  $M = 3$  modes.

should be identical and this common per-RB transmit energy  $\rho$  must be chosen from a finite set  $\mathcal{P}$ . We justifiably assume that  $\mathcal{P}$  has only strictly positive entries. Indeed, this constraint models a common way of operating LTE networks. In the current LTE deployments (that conform to LTE Release 8), the users attached to a cell are first conveyed a cell-specific reference signal (CRS) power spectral density using an integer which spans  $\{-60, \dots, 50\}$ , where each value conveys power in dBm per 15 kHz [84]. The transmit energy expended on each RB that is used for data transmission to a scheduled user is a scalar times this CRS power spectral density. Typical operation employs a cell-specific scalar that is identical for all users attached to that cell, thereby resulting in an identical energy spent on each used RB. We note here that the LTE downlink aims to extract frequency selective scheduling gains (enabled due to OFDMA). It exploits the observation that further optimizing transmit powers in the frequency domain (in conjunction with frequency selective scheduling) provides marginal gains at the cost of increased signaling support and complexity. In addition, such power optimization can have a detrimental effect on ACK/NACK based link adaptation. In



this context, we recall a similar observation made in [85] that constant power allocation (over a set of “good” channels) is close to the optimal power allocation over those channels (determined via waterfilling).

**[C5]:** The per-RB energy,  $\rho$ , together with the set of used RBs should not violate the energy harvesting constraints. The energy causality requires that the energy used till the end of any subframe in the block doesn’t exceed the energy available by that subframe. Consequently, for any given energy per-RB,  $\rho$ , and the energy harvested per-subframe  $\ell$ ,  $E_\ell$ , we can first determine  $L$  RB cardinality limits  $J_1, \dots, J_L$  which define a system of linear inequalities. In particular, the number of assigned RBs in the first sub-frame should not exceed

$$J_1 = \left\lfloor \frac{E_0 + E_1 - \vartheta}{\gamma \rho} \right\rfloor,$$

where  $E_0$  denotes the energy stored in the battery at the start of the block,  $\vartheta \geq 0$  represents the circuit energy consumed over each subframe and  $\gamma \geq 1$  is the inverse of the power amplifier efficiency. Further, the number of assigned RBs in the first and second sub-frames together should not exceed

$$J_1 + J_2 = \left\lfloor \frac{E_0 + E_1 + E_2 - 2\vartheta}{\gamma \rho} \right\rfloor,$$

and so on. For simplicity, we assume that the incoming energy in each subframe  $\ell$  can at least provide the circuit energy, i.e.,  $E_\ell \geq \vartheta \forall \ell$ . Alternatively, we can also suppose that the circuit energy is supplied by another non-renewable source in which case we have

$$J_1 = \left\lfloor \frac{E_0 + E_1}{\gamma \rho} \right\rfloor,$$

and

$$J_1 + J_2 = \left\lfloor \frac{E_0 + E_1 + E_2}{\gamma \rho} \right\rfloor,$$

and so on. Another variation where RBs are used in a block only after guaranteeing that the energy available at the start of the next one exceeds the required circuit energy, can also be incorporated. In each case, these linear inequalities arise due to increasing

renewable energy that is available as we traverse across the subframes of the block. A feasible resource allocation is depicted in Fig. 7.1.

### 7.1.2 Objective Functions

We consider the following two energy efficiency metrics, which have been defined to conform to identical energy per RB constraint (C4).

**Global Energy Efficiency (GEE):** Ratio of the weighted sum rate (computed per block) and the corresponding energy consumption (Joules per block)

$$\text{GEE}(\rho, J, \mathcal{R}) = \frac{\sum_{n \in \mathcal{R}} \psi_n r_n}{L\vartheta + J\gamma\rho}, \quad (7.1)$$

where  $\mathcal{R} \subseteq \mathcal{N}$  denotes the set of RBs chosen for data transmission over  $L$  subframes.  $J$  denotes the number of used RBs (i.e., cardinality of the set  $\mathcal{R}$ ) and  $r_n$  denotes the rate (throughput) achieved on RB  $n \in \mathcal{R}$ .  $\psi_n, \forall n \in \mathcal{N}$  denote the weights (or priorities) specified for all the available RBs.

**Weighted Sum of Energy Efficiencies (WSEE):**

$$\text{WSEE}(\rho, J, \mathcal{R}) = \sum_{n \in \mathcal{R}} \psi_n \frac{r_n}{\delta_n + \gamma\rho}, \quad (7.2)$$

where  $\delta_n > 0$  is the given circuit energy consumed corresponding to RB  $n$ . A default choice could be to set each  $\delta_n$  as the total circuit energy  $L\vartheta$  amortized over all  $NL$  RBs in the block, i.e.,  $\delta_n = \vartheta/N$  for all  $n$ .

Each one of the two metrics will be optimized (in each scheduling block) over the choice of the per-RB energy  $\rho$ , the number of used RBs,  $J$ , the set  $\mathcal{R}$  of cardinality  $J$  and the choice of users (with their respective modes and bit loading) scheduled on those RBs, subject to the set of constraints C1-to-C5 described above. Our first observation is:

**Lemma 7.1** *Without loss of optimality, we can replace each  $J_\ell$  by  $J'_\ell, \forall \ell = 1, \dots, L$ , where the integers  $J'_\ell : J'_\ell \leq N, \forall \ell$  are recursively defined as  $J'_\ell = \min\{N, J_\ell + \sum_{q=1}^{\ell-1} (J_q - J'_q)\}, \forall \ell$ .*

Indeed, any set of RBs  $\mathcal{R} \subseteq \mathcal{N}$  that is feasible under the linear inequalities defined by  $\{J_\ell\}_{\ell=1}^L$ , i.e., satisfies  $|\mathcal{R} \cap \{\cup_{j=1}^\ell \mathcal{N}_j\}| \leq \sum_{j=1}^\ell J_j$ ,  $\forall \ell = 1, \dots, L$ , also remains feasible under the linear inequalities defined by  $\{J'_\ell\}_{\ell=1}^L$  and vice versa. Consequently, with a slight abuse of notation, henceforth we use  $J_\ell$  to denote  $J'_\ell$ , so that  $J_\ell \leq N$ ,  $\forall \ell$ .

We now define the two problems of interest as

$$\max\{\text{GEE}(\rho, J, \mathcal{R})\} \text{ s.t. C1, C2, C3, C4 \& C5, \quad (\text{PI})$$

$$\max\{\text{WSEE}(\rho, J, \mathcal{R})\} \text{ s.t. C1, C2, C3, C4 \& C5, \quad (\text{PII})$$

In order to solve these two problems, we let  $x_{u,n}^m$  be the indicator variable which is one if user  $u$  is scheduled on RB  $n$  with transmission mode  $m$ , and zero otherwise, and define a key sub-problem of *maximizing the weighted sum rate* for any given  $\rho$  and  $J$ .

$$\max_{\{x_{u,n}^m, b_{u,n}^m\}} \sum_{u=1}^K \sum_{\ell=1}^L \sum_{n \in \mathcal{N}_\ell} \sum_{m=1}^M \psi'_{u,n} x_{u,n}^m b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho)) \quad (7.3a)$$

$$\text{subject to } \sum_{m=1}^M \sum_{u=1}^K x_{u,n}^m \leq 1, \quad n \in \mathcal{N} \quad (7.3b)$$

$$\sum_{u=1}^K \sum_{m=1}^M \max_{n \in \mathcal{N}} x_{u,n}^m \leq \bar{K} \quad (7.3c)$$

$$\sum_{m=1}^M \max_{n \in \mathcal{N}} x_{u,n}^m \leq 1, \quad 1 \leq u \leq K \quad (7.3d)$$

$$\sum_{n \in \mathcal{N}_1 \cup \dots \cup \mathcal{N}_\ell} \sum_{u=1}^K \sum_{m=1}^M x_{u,n}^m \leq \min \left\{ \sum_{q=1}^\ell J_q, J \right\}, \forall \ell \quad (7.3e)$$

$$\sum_{m=1}^M \sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u, \quad 1 \leq u \leq K \quad (7.3f)$$

$$x_{u,n}^m \in \{0, 1\} \text{ \& } b_{u,n}^m \in \mathcal{B}^m, \forall u, m, n. \quad (7.3g)$$

where  $\{\psi'_{u,n}\}_{u \in \mathcal{U}, n \in \mathcal{N}}$  are given weights, and  $b_{u,n}^m$  denotes the number of information bits allocated on RB  $n$  when user  $u$  is assigned to that RB with transmission mode  $m$ . The set of constraints (7.3b) and (7.3c) describe the OFDMA orthogonality constraint and the user limit constraint, respectively (constraint C1). The set of constraints (7.3d) stipulates that each user can only have one transmission mode across all RBs

allocated to it (constraint C2). Similarly, the constraint (7.3e) dictates that the total number of all occupied RBs should be no greater than  $J$  while satisfying the energy causality constraints (C5). The set of constraints (7.3f) and (7.3g) enforce the queue size limit for each user and permissible bit loadings (constraint C3). We note that (7.3) is a more practical and fully combinatorial analogue of the problem considered in [80]. We can verify that (7.3) subsumes the problem in [57] if we consider the 0 – 1 throughput model (6.3) with just one available mode and no linear causality constraints or cardinality bound on the set of used RBs. So, (7.3) is at least as hard as the one posed in [57]. Invoking the hardness result established for the latter problem in [57], we can deduce that (7.3) is NP-hard. Notice also that (7.3) is a sub-problem of both (PI) and (PII), obtained for any choice of  $\rho$  and  $J$ , upon setting  $\psi'_{u,n} = \psi_n/(L\vartheta + \gamma\rho J)$  and  $\psi'_{u,n} = \psi_n/(\delta_n + \gamma\rho)$ ,  $\forall n$ , respectively. These observations can be used to show that both (PI) and (PII) are intractable. Moreover, we also can assert the following result.

**Proposition 7.1** *For any constant  $\alpha \in (0, 1]$ , an  $\alpha$ –approximation algorithm for (7.3) can be used to design  $\alpha$ –approximation algorithms for each one of the problems (PI) and (PII).*

**Proof:** We only prove the result for GEE since the one for WSEE is straightforward. Suppose that we have an  $\alpha$ –approximation algorithm for (7.3) that for any given choice of  $\rho, J$ , and other input parameters provides a feasible solution (satisfying constraints C1-to-C3 and C5) which yields an objective value no less than  $\alpha$  times the optimal objective value for that input. Then, for each  $\rho \in \mathcal{P}$  and each  $J \in \{1, \dots, \sum_{\ell=1}^L J_\ell\}$ , let us invoke the algorithm at hand (with weights  $\psi'_{u,n} = \psi_n/(L\vartheta + \gamma\rho J)$ ) to obtain a feasible solution and let  $\hat{O}(\rho, J)$  denote the corresponding weighted sum rate. We can deduce that  $\hat{O}(\rho, J) \geq \alpha \max\{\text{GEE}(\rho, J, \mathcal{R})\}$ , where the latter maximization is subject to constraints C1-C3 and C5. Thus, by selecting the best among solutions obtained after considering all finitely many feasible choices of  $\rho, J$  (which can be no greater than  $NL|\mathcal{P}|$ ), we obtain one yielding a GEE no less than  $\alpha$  times the optimal GEE.  $\square$

As a result we focus on designing a constant-factor approximation algorithm for (7.3),

which as proved above yields constant-factor approximation algorithms for both (PI) and (PII). Before that we demonstrate one useful exception when the optimal GEE can be efficiently determined. Towards this end, for each mode  $m$  and user  $u$ , we define a set  $\mathcal{B}_u^m = \{b \in \mathcal{B}^m : b \leq Q_u\}$ . Without loss of generality, we can assume that the set  $\mathcal{B}_u^m$  has at-least one strictly positive member (otherwise we can simply remove mode  $m$  as a candidate mode for user  $u$ ). Further, for each RB  $n$  we let

$$\bar{r}_{u,n}^m(\rho) = \max_{b_n \in \mathcal{B}_u^m} \{b_n(1 - p_{u,n}^m(b_n, \rho))\},$$

and recall that the set  $\mathcal{P}$  is assumed to contain only strictly positive entries.

**Proposition 7.2** *Suppose that for each  $\rho \in \mathcal{P}$  there exists*

$$(u^*, m^*, n^*) = \arg \max_{u \in \mathcal{U}, m \in \mathcal{M}, n \in \mathcal{N}} \{\psi_n \bar{r}_{u,n}^m(\rho)\}.$$

*If*

$$\frac{\psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho)}{L\vartheta + \gamma\rho} \geq \max_{\substack{u \in \mathcal{U}, n \in \mathcal{N} : (u, n) \neq (u^*, n^*) \\ \& m \in \mathcal{M}}} \left\{ \frac{\psi_n \bar{r}_{u,n}^m(\rho)}{\gamma\rho} \right\}, \quad (7.4)$$

*then, an optimal solution to (PI) is given by*

$$\max_{\rho \in \mathcal{P}} \max_{u \in \mathcal{U}, m \in \mathcal{M}, n \in \mathcal{N}} \left\{ \frac{\psi_n \bar{r}_{u,n}^m(\rho)}{L\vartheta + \gamma\rho} \right\}. \quad (7.5)$$

**Proof:** The simple but key observation that supports this proposition is that for any non-negative scalars  $a, b, c, d$ , such that  $b \neq 0, d \neq 0$ ,

$$\frac{a}{b} \geq \frac{c}{d} \Leftrightarrow \frac{a}{b} \geq \frac{a+c}{b+d} \geq \frac{c}{d}. \quad (7.6)$$

Recalling (7.1) let  $\hat{\mathcal{R}}, \hat{J}$  denote the set of RBs and its cardinality that is optimal with respect to the GEE at some  $\rho$ , and let  $\{\hat{R}_n\}_{n \in \hat{\mathcal{R}}}$  denote the weighted rates achieved on those RBs. Further, recall that each RB can be assigned to at most one user with one mode and expand  $\hat{\mathcal{R}} = \{n_1, \dots, n_j\}$  where  $\hat{R}_{n_1} \geq \hat{R}_{n_2} \geq \dots \geq \hat{R}_{n_j}$ . Assume (7.4)

holds at that  $\rho$  (where we note that  $(u^*, m^*, n^*)$  can depend on  $\rho$ ). Suppose first that user  $u^*$  is present in (or scheduled under) the optimal solution at hand. In this case, let  $n_\ell \in \{n_1, \dots, n_{\hat{J}}\}$  be an RB such that  $n_\ell = n^*$  if user  $u^*$  has been assigned RB  $n^*$  in that optimal solution, otherwise  $n_\ell$  can be chosen to be any RB assigned to user  $u^*$  in that optimal solution. Then, invoking (7.4) with (7.6) we see that

$$\frac{\psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho)}{L\vartheta + \gamma\rho} \geq \frac{\hat{R}_n + \psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho)}{L\vartheta + 2\gamma\rho} \geq \frac{\hat{R}_n}{\gamma\rho}$$

holds for all  $n \in \hat{\mathcal{R}} \setminus \{n_\ell\}$ , whereas  $\psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho) \geq \hat{R}_{n_\ell}$ . Then, since

$$\frac{\hat{R}_{n_1}}{\gamma\rho} \geq \frac{\hat{R}_{n_2}}{\gamma\rho} \geq \dots \geq \frac{\hat{R}_{n_{\hat{J}}}}{\gamma\rho},$$

we can repeatedly invoke (7.6) to obtain

$$\frac{\psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho)}{L\vartheta + \gamma\rho} \geq \frac{\psi_{n^*} \bar{r}_{u^*, n^*}^{m^*}(\rho) + \sum_{n \in \hat{\mathcal{R}} \setminus \{n_\ell\}} \hat{R}_n}{L\vartheta + \gamma\rho + (\hat{J} - 1)\gamma\rho} \geq \frac{\sum_{n \in \hat{\mathcal{R}}} \hat{R}_n}{L\vartheta + \hat{J}\gamma\rho},$$

which settles the result. The other case where the user  $u^*$  is not present in the solution at hand can be proved analogously.  $\square$

Notice that since the set  $\mathcal{P}$  is assumed to contain only strictly positive entries, the GEE is well defined even when  $\vartheta = 0$  and indeed the condition in (7.4) is always satisfied when  $\vartheta = 0$ . Proposition 7.2 implies that for all small enough  $\vartheta$ , a GEE optimal solution involves using only one RB and hence can be efficiently determined.

## 7.2 A Constant Factor Approximation Algorithm

We next proceed to re-formulate (7.3). Towards this end, we define a ground set containing all possible 3-tuples  $\Psi = \{(u, m, \mathbf{b}_u^m)\}$ , where for each such 3-tuple (or element)  $\underline{e} = (u, m, \mathbf{b}_u^m)$ ,  $u$  denotes the user and  $m$  denotes the mode and  $\mathbf{b}_u^m = [b_{u,1}^m, \dots, b_{u,NL}^m]^T \in \mathbb{R}_+^{NL}$  is an  $NL$  length bit loading vector such that on each RB  $n \in \mathcal{N}$ , the bit loading is permissible,  $b_{u,n}^m \in \mathcal{B}^m$  and across all RBs the buffer size constraint is satisfied, i.e.,  $\sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u$ . Furthermore, we define a family of subsets

of  $\Psi$ , denoted by  $\underline{\mathcal{I}}$ , as

$$\underline{\mathcal{I}} = \left\{ \mathcal{A} \subseteq \Psi : |\mathcal{A}| \leq \bar{K}, \quad \sum_{(u', m', \mathbf{b}_{u'}^{m'}) \in \mathcal{A}} \mathbf{1}(u' = u) \leq 1 \quad \forall u \in \mathcal{U} \right\}.$$

In words, any subset of 3-tuples from  $\Psi$  in which each user appears at-most once and whose cardinality does not exceed the user limit  $\bar{K}$  is a member of  $\underline{\mathcal{I}}$ . Then, we define a normalized non-negative set function  $h : 2^\Psi \rightarrow \mathbb{R}_+$ , such that  $h(\emptyset) = 0$  and for all other subsets  $\mathcal{A} \subseteq \Psi$

$$\begin{aligned} h(\mathcal{A}) &= \max_{\substack{z_n \in \{0,1\} \\ \forall n}} \left\{ \sum_{n \in \mathcal{N}} z_n \max_{(u, m, \mathbf{b}_u^m) \in \mathcal{A}} \{ \psi'_{u,n} b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho)) \} \right\}, \\ \text{s.t.} \quad \sum_{n \in \cup_{j=1}^{\ell} \mathcal{N}_j} z_n &\leq \min \left\{ J, \sum_{q=1}^{\ell} J_q \right\}, \quad \ell = 1, \dots, L. \end{aligned} \quad (7.7)$$

We can now reformulate the problem in (7.3) as

$$\begin{aligned} &\max_{\mathcal{A} \subseteq \Psi} \{ h(\mathcal{A}) \} \\ \text{s.t.} \quad &\mathcal{A} \in \underline{\mathcal{I}}. \end{aligned} \quad (7.8)$$

Some comments on this reformulation are in order. First, from the definition of  $\underline{\mathcal{I}}$  it follows that by restricting  $\mathcal{A} \in \underline{\mathcal{I}}$  we have ensured that each user is selected at-most once with one mode and that the associated bit loading is feasible, thereby meeting the per-user mode and bit loading constraints of (7.3). Then, the definition of the set function  $h(\cdot)$  ensures that each RB is implicitly assigned to at-most one user (since only the weighted throughput of at-most one user is chosen via the inner  $\max(\cdot)$  function) and that the linear causality constraints are imposed in determining the weighted sum rate (via the indicator variables  $\{z_n\}$ ). Consequently, any feasible solution to (7.8) maps to a feasible one for (7.3) yielding the same objective value and vice versa. Then, we have the following theorem which is our main result.

**Theorem 7.1** *The set function  $h(\cdot)$  is a normalized non-decreasing submodular set*

function and  $(\Psi, \underline{\mathcal{I}})$  is a matroid.

**Proof:** The fact that  $(\Psi, \underline{\mathcal{I}})$  is a matroid follows upon verifying the exchange property stated in the appendix of Chapter 6. Moreover, it can also be readily verified that the set function  $h(\cdot)$  is a normalized non-decreasing set function. Then, to prove the submodularity of this set function we first establish that

$$h(\mathcal{A} \cup \underline{e}') - h(\mathcal{A}) \geq h(\mathcal{B} \cup \underline{e}') - h(\mathcal{B}), \quad (7.9)$$

for any two subsets  $\mathcal{A}, \mathcal{B} \subseteq \Psi$  with  $\mathcal{A} \subseteq \mathcal{B}$  and any 3-tuple  $\underline{e}' = (u', m', \mathbf{b}_{u'}^{m'}) \in \Psi$  whose bit loading vector  $\mathbf{b}_{u'}^{m'}$  satisfies  $b_{u',t}^{m'} \geq 0$  for any one RB  $t \in \mathcal{N}$ , whereas  $b_{u',n}^{m'} = 0 \forall n \in \mathcal{N}, n \neq t$ . Let  $\mathcal{E} \subseteq \Psi$  denote the set of all such 3-tuples whose bit loading vectors have a positive entry in at-most one RB. We first define an  $NL$  length vector  $\Delta^{\mathcal{A}}$ , where

$$\Delta_n^{\mathcal{A}} = \max_{(u,m,\mathbf{b}_u^m) \in \mathcal{A}} \{\psi'_{u,n} b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho))\} \quad n \in \mathcal{N}.$$

Then, we define  $L$  sets  $\mathcal{R}_\ell^{\mathcal{A}}$ ,  $\ell = 1, \dots, L$  in a recursive manner as follows. The set  $\mathcal{R}_1^{\mathcal{A}}$  is defined as the set containing the  $\min\{J, J_1\}$  RBs corresponding to the  $\min\{J, J_1\}$  largest members of  $\{\Delta_n^{\mathcal{A}}\}_{n \in \mathcal{N}_1}$ . Each subsequent set  $\mathcal{R}_\ell^{\mathcal{A}}$ ,  $\ell = 2, \dots, L$  is defined as the set containing the  $\min\{J, J_1 + \dots + J_\ell\}$  RBs corresponding to the  $\min\{J, J_1 + \dots + J_\ell\}$  largest members of  $\{\Delta_n^{\mathcal{A}}\}_{n \in \mathcal{N}_\ell \cup \mathcal{R}_{\ell-1}^{\mathcal{A}}}$ . Thus, we have the following telescoping relations,

$$\mathcal{R}_L^{\mathcal{A}} \cap \mathcal{N}_\ell \subseteq \mathcal{R}_{L-1}^{\mathcal{A}} \cap \mathcal{N}_\ell \subseteq \dots \subseteq \mathcal{R}_\ell^{\mathcal{A}} \cap \mathcal{N}_\ell, \quad \ell = 1, \dots, L. \quad (7.10)$$

With this definition, note that each  $\mathcal{R}_\ell^{\mathcal{A}}, \forall \ell$  is an optimal set of RBs (for the given  $\mathcal{A}$ ) chosen from  $\mathcal{N}_1 \cup \dots \cup \mathcal{N}_\ell$  that maximizes the weighted sum rate subject to the first  $\ell$  causality constraints. Hence, we can compute  $h(\mathcal{A})$  using the objective in (7.7) after setting indicator variables  $\{z_n\}$  to be one for all RBs in  $\mathcal{R}_L^{\mathcal{A}}$  and zero for all other RBs. Next, we define the weighted rate barrier at each RB  $n$ ,  $V_n^{\mathcal{A}}, n \in \mathcal{N}_\ell, \forall \ell$ , as

$$V_n^{\mathcal{A}} = \max \left\{ \Delta_n^{\mathcal{A}}, \max \{ V_{\min, \ell}^{\mathcal{A}}, V_{\min, \ell+1}^{\mathcal{A}}, \dots, V_{\min, L}^{\mathcal{A}} \} \right\}, \quad (7.11)$$



where  $V_{\min,\ell}^{\mathcal{A}} = \min_{n \in \mathcal{R}_\ell^{\mathcal{A}}} \{\Delta_n^{\mathcal{A}}\}$ ,  $\forall \ell : 1 \leq \ell \leq L$ . The operational meaning of the weighted rate barrier at RB  $n$ ,  $V_n^{\mathcal{A}}$ , is that if we augment  $\mathcal{A}$  by adding any 3-tuple  $\underline{e} = (u, m, \mathbf{b}_u^m) \in \mathcal{E} : b_{u,n}^m > 0$ , then

$$h(\mathcal{A} \cup \underline{e}) - h(\mathcal{A}) = [\psi'_{u,n} b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho)) - V_n^{\mathcal{A}}]^+.$$

In other words, the weighted throughput offered by the added 3-tuple on RB  $n$ ,  $\psi'_{u,n} b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho))$ , must exceed the barrier  $V_n^{\mathcal{A}}$  to improve the objective.

In an analogous manner, we define the vector  $\Delta^{\mathcal{B}}$ , the sets  $\mathcal{R}_\ell^{\mathcal{B}}$ ,  $\ell : 1 \leq \ell \leq L$ , and the weighted rate barrier at each RB  $n$ ,  $V_n^{\mathcal{B}}$ ,  $n \in \mathcal{N}$ . The following deduction is straightforward

$$\Delta_n^{\mathcal{A}} \leq \Delta_n^{\mathcal{B}}, \quad n \in \mathcal{N}, \quad (7.12)$$

whereas

$$V_{\min,\ell}^{\mathcal{A}} \leq V_{\min,\ell}^{\mathcal{B}}, \quad 1 \leq \ell \leq L,$$

follows from (7.12) and the definitions of sets  $\mathcal{R}_\ell^{\mathcal{A}}, \mathcal{R}_\ell^{\mathcal{B}}, l$  after some algebra. Combining these two facts, we obtain that

$$V_n^{\mathcal{A}} \leq V_n^{\mathcal{B}}, \quad n \in \mathcal{N}. \quad (7.13)$$

Next, upon defining  $G = \psi'_{u',t} b_{u',t}^{m'} (1 - p_{u',t}^{m'}(b_{u',t}^{m'}, \rho))$  we can see that  $h(\mathcal{A} \cup \underline{e}') - h(\mathcal{A}) = [G - V_t^{\mathcal{A}}]^+$  and  $h(\mathcal{B} \cup \underline{e}') - h(\mathcal{B}) = [G - V_t^{\mathcal{B}}]^+$ . Invoking (7.13) it can now be verified that (7.9) holds true. We will leverage this result to show that (7.9) holds for any 3-tuple  $\underline{e}' = (u', m', \mathbf{b}_{u'}^{m'}) \in \Omega \setminus \mathcal{B}$ , without the restriction that  $\underline{e}' \in \mathcal{E}$ , thereby proving the theorem. For convenience, we adopt the notation that for all  $n = 1, \dots, NL$ ,  $\mathbf{b}_{u',(n)}^{m'}$  denotes the  $NL$  length vector formed by retaining the first  $n$  components of  $\mathbf{b}_{u'}^{m'}$  and setting the remaining ones to zero, i.e.,

$$\mathbf{b}_{u',(n)}^{m'} = [b_{u',1}^{m'}, \dots, b_{u',n}^{m'}, 0, \dots, 0].$$

Similarly, we let  $\mathbf{b}_{u',\bar{n}}^{m'}$  denote the  $NL$  length vector formed by retaining only the  $n^{th}$  component of  $\mathbf{b}_{u'}^{m'}$  and setting all the other ones to zero, i.e.,

$$\mathbf{b}_{u',\bar{n}}^{m'} = [0, \dots, 0, b_{u',n}^{m'}, 0, \dots, 0].$$

Considering the difference  $h(\mathcal{B} \cup \underline{e}') - h(\mathcal{B})$ , we expand it as

$$\begin{aligned} h(\mathcal{B} \cup \underline{e}') - h(\mathcal{B}) &= h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(1)}^{m'})) - h(\mathcal{B}) + \\ &\sum_{n=2}^{NL} \left( h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n)}^{m'})) - h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \right) \end{aligned} \quad (7.14)$$

From the result we have proved, we see that since  $(u', m', \mathbf{b}_{u',(1)}^{m'}) \in \mathcal{E}$ ,

$$h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(1)}^{m'})) - h(\mathcal{B}) \leq h(\mathcal{A} \cup (u', m', \mathbf{b}_{u',(1)}^{m'})) - h(\mathcal{A}).$$

Notice then from the definition of the set function  $h(\cdot)$  (7.7), we have that for any set  $\mathcal{B} \subseteq \Omega$ ,

$$\begin{aligned} &h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n)}^{m'})) - h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \\ &= h((\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \cup (u', m', \mathbf{b}_{u',\bar{n}}^{m'})) \\ &\quad - h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})). \end{aligned}$$

Noting that since each 3-tuple  $(u', m', \mathbf{b}_{u',\bar{n}}^{m'}) \in \mathcal{E}$ , it allows us to again invoke the result proved earlier to deduce this fact

$$\begin{aligned} &h((\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \cup (u', m', \mathbf{b}_{u',\bar{n}}^{m'})) \\ &\quad - h(\mathcal{B} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \\ &\leq h((\mathcal{A} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \cup (u', m', \mathbf{b}_{u',\bar{n}}^{m'})) \\ &\quad - h(\mathcal{A} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})) \\ &= h(\mathcal{A} \cup (u', m', \mathbf{b}_{u',(n)}^{m'})) - h(\mathcal{A} \cup (u', m', \mathbf{b}_{u',(n-1)}^{m'})). \end{aligned}$$

Combining this fact and (7.14), we get the desired result.  $\square$

From Theorem 7.1 we can infer that the reformulated problem in (7.8) is one of maximizing a submodular objective subject to one matroid constraint. Thus, we can adapt the classical greedy method [71] for the latter optimization problem. However, the main hurdle we need to overcome is that of selecting the locally optimal 3-tuple (given a set of selected 3-tuples). In particular, the sub-problem we have to solve at each iteration, given a set  $\hat{\mathcal{G}}$  of 3-tuples selected so far, for each un-selected user  $u$  and mode  $m$  can be posed as

$$\max_{\substack{\mathbf{b}_u^m \in \mathbb{R}_+^{NL} \\ b_{u,n}^m \in \mathcal{B}^m \forall n \in \mathcal{N} \text{ \& } \sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u}} \{h(\hat{\mathcal{G}} \cup (u, m, \mathbf{b}_u^m)) - h(\hat{\mathcal{G}})\}. \quad (7.15)$$

An important observation that follows from the submodularity of  $h(\cdot)$  is the following.

**Lemma 7.2** *The problem in (7.15) is the maximization of a normalized non-decreasing submodular set function subject to one knapsack constraint.*

Thus, (7.15) can itself be solved approximately (with a constant-factor  $\eta = 1 - \frac{1}{\sqrt{e}}$  guarantee) using an enhanced greedy method [75]. Notice that the knapsack constraint becomes vacuous for the full buffer traffic model and in this case the enhanced greedy method returns the optimal solution. Our proposed algorithm is an adaptation of the classical greedy algorithm that at each iteration invokes the enhanced greedy method, to the particular problem at hand (7.8). Then, invoking the approximation guarantee derived for the classical greedy method when used to maximize a submodular set function under a matroid constraint, with an approximately locally optimal choice at each step, [86], we obtain the following.

**Theorem 7.2** *Let  $\mathcal{O}^{opt}$  denote an optimal solution to (7.8) and let  $\mathcal{O}$  denote the one yielded by our proposed Algorithm. Then,*

$$h(\mathcal{O}) \geq \frac{\eta h(\mathcal{O}^{opt})}{\eta + 1}, \quad (7.16)$$

where the constant  $\eta = 1 - \frac{1}{\sqrt{e}}$ . Further, in the special case of backlogged (full buffer) traffic model  $\eta = 1$ .

We now provide a detailed description of our proposed algorithm. Towards this end, recalling the definition of the subset  $\mathcal{E} \subseteq \Psi$  as the set of all 3-tuples whose bit loading vectors have a positive entry in at-most one RB, we introduce our adaptation of the classical greedy method in Algorithm 3, which repeatedly invokes an adaptation of the enhanced greedy method, described in Algorithm 4. On perusing Algorithms 3 and 4 we see that the key step that needs to be efficiently solved is (7.21) (or (7.22)). This can be done as follows. Let  $\mathcal{A} = \hat{\mathcal{G}} \cup \hat{\mathcal{e}}$  and consider any  $\underline{e} = (u, m, \mathbf{b}_u^m) \in \mathcal{E}$  for which  $b_{u,n}^m > 0$ . Then, from the proof of Theorem 7.1 we can deduce that

$$h(\hat{\mathcal{G}} \cup \hat{\mathcal{e}} \cup \underline{e}) - h(\hat{\mathcal{G}} \cup \hat{\mathcal{e}}) = [\psi'_{u,n} b_{u,n}^m (1 - p(b_{u,n}^m, \rho)) - V_n^{\mathcal{A}}]^+.$$

Therefore given any subset  $\mathcal{A} \subseteq \Psi$ , updating the weighted rate barrier upon inclusion of a new 3-tuple  $\underline{e}$ , i.e., determining  $V_n^{\mathcal{A} \cup \underline{e}}$ ,  $\forall n \in \mathcal{N}$ , is the key hurdle we have to surmount. This is achieved by the following result.

**Proposition 7.3** *For any given subset  $\mathcal{A} \subseteq \Psi$  and its corresponding optimal RB set  $\mathcal{R}_L^{\mathcal{A}}$ , define the set of bottleneck subframes in the scheduling block as*

$$\mathcal{S}^{\mathcal{A}} = \left\{ \ell \in \{1, \dots, L\} : |\mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=1}^{\ell} \mathcal{N}_j)| = \min \left\{ J, \sum_{j=1}^{\ell} J_j \right\} \right\}. \quad (7.17)$$

*Then, the weighted rate barrier  $V_n^{\mathcal{A}}$  for any  $n \in \mathcal{N}_\ell, \ell = 1, \dots, L$  can also be written as*

$$V_n^{\mathcal{A}} = \max \left\{ \Delta_n^{\mathcal{A}}, \min_{s \in \mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j)} \{\Delta_s^{\mathcal{A}}\} \right\}, \quad (7.18)$$

*where  $\hat{\ell} = \min\{j \in \mathcal{S}^{\mathcal{A}} : j \geq \ell\}$ . Furthermore, upon inclusion of any new 3-tuple  $\underline{e} = (u, m, \mathbf{b}_u^m) \in \mathcal{E}$ , such that  $b_{u,n}^m > 0$  for that  $n \in \mathcal{N}_\ell$ , the set  $\mathcal{R}_L^{\mathcal{A} \cup \underline{e}}$  can be determined using*

$$\mathcal{R}_L^{\mathcal{A} \cup \underline{e}} = \begin{cases} \mathcal{R}_L^{\mathcal{A}} & D \leq 0 \text{ or } n \in \mathcal{R}_L^{\mathcal{A}}, \\ (\mathcal{R}_L^{\mathcal{A}} \setminus \{\check{s}\}) \cup \{n\} & \text{else,} \end{cases} \quad (7.19)$$

where

$$D = \psi'_{u,n} b_{u,n}^m (1 - p(b_{u,n}^m, \rho)) - V_n^{\mathcal{A}},$$

and

$$\check{s} = \arg \min_{s \in \mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j)} \{\Delta_s^{\mathcal{A}}\},$$

and ties can be broken arbitrarily.

**Proof:** We first show that the alternate expression given for the weighted rate barrier,  $V_n^{\mathcal{A}}$ ,  $n \in \mathcal{N}_{\ell}$ , in (7.18) is indeed equivalent to the one in (7.11). Note first that the last subframe must always be a bottleneck subframe, i.e.,  $L \in \mathcal{S}^{\mathcal{A}}$ ,  $\forall \mathcal{A}$ . Consider a 3-tuple (or element)  $\underline{e} = (u, m, \mathbf{b}_u^m) \in \mathcal{E}$  whose bit vector has a positive entry on the RB  $n$  of interest ( $b_{u,n}^m > 0$ ) and let  $R = \psi'_{u,n} b_{u,n}^m (1 - p(b_{u,n}^m, \rho))$  denote the weighted rate of element  $\underline{e}$  on RB  $n$ . Then, to ensure an improvement in the utility upon adding  $\underline{e}$  to  $\mathcal{A}$ , a trivial necessary condition is  $R > \Delta_n^{\mathcal{A}}$  so that  $V_n^{\mathcal{A}} \geq \Delta_n^{\mathcal{A}}$ . Next, let  $\hat{\ell}$  be the first bottleneck subframe that is either at or after the subframe  $\ell$  which contains the RB  $n$  of interest. By definition of a bottleneck subframe, the optimal set  $\mathcal{R}_L^{\mathcal{A}}$  includes the maximum possible number of RBs from subframes  $1, \dots, \hat{\ell}$ . This implies that the best weighted rates obtained over RBs in subframes  $\hat{\ell} + 1, \dots, L$  under the set  $\mathcal{A}$ , i.e.,  $\{\Delta_s^{\mathcal{A}}\}, \forall s \in \cup_{j=\hat{\ell}+1}^L \mathcal{N}_j$ , are not collectively large enough to ensure that a higher weighted sum rate utility can be obtained by assigning fewer than  $\min\{J, \sum_{j=1}^{\hat{\ell}} J_j\}$  RBs to the first  $\hat{\ell}$  subframes. Consequently, we can deduce that the telescoping relation in (7.10) can be further refined to

$$\mathcal{R}_L^{\mathcal{A}} \cap \mathcal{N}_j = \mathcal{R}_{L-1}^{\mathcal{A}} \cap \mathcal{N}_j = \dots = \mathcal{R}_{\hat{\ell}}^{\mathcal{A}} \cap \mathcal{N}_j, \forall j = 1, \dots, \hat{\ell}, \quad (7.20)$$

which also means that

$$\mathcal{R}_{\hat{\ell}}^{\mathcal{A}} = \mathcal{R}_L^{\mathcal{A}} \cap (\cup_{k=1}^{\hat{\ell}} \mathcal{N}_k) \subseteq \mathcal{R}_j^{\mathcal{A}}, \quad j = \hat{\ell} + 1, \dots, L,$$

so that

$$V_{\min, \hat{\ell}}^{\mathcal{A}} \geq V_{\min, j}^{\mathcal{A}}, \quad j = \hat{\ell} + 1, \dots, L.$$

This fact upon invoking (7.11) proves the equivalence of the expressions in (7.18) and (7.11) whenever  $\hat{\ell} = \ell$ . Then, suppose that  $\hat{\ell} > \ell$ . Here, it can be seen that the best weighted rate on each RB  $s \in \mathcal{R}_{\ell'}^{\mathcal{A}} \setminus \mathcal{R}_{\hat{\ell}}^{\mathcal{A}}, \Delta_s^{\mathcal{A}}$ , is no greater than the best weighted rate on any RB  $t \in \mathcal{R}_{\hat{\ell}}^{\mathcal{A}} \setminus \mathcal{R}_{\ell'}^{\mathcal{A}}, \Delta_t^{\mathcal{A}}$ , for all  $\ell' = \ell, \dots, \hat{\ell} - 1$ . This is because otherwise by swapping RB  $t$  in  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A}} \setminus \mathcal{R}_{\ell'}^{\mathcal{A}}$  by RB  $s$  in  $\mathcal{R}_{\ell'}^{\mathcal{A}} \setminus \mathcal{R}_{\hat{\ell}}^{\mathcal{A}}$ , we can obtain a higher weighted sum rate while retaining feasibility. In other words, such a swap will yield a feasible set of RBs in  $\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j$  because of (7.10) and the fact that subframes  $\ell, \dots, \hat{\ell} - 1$  are not bottleneck subframes, and this feasible set will provide a higher weighted sum rate (under set  $\mathcal{A}$ ) which contradicts the weighted sum rate optimality of  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A}}$ . Thus, we have the additional result  $V_{\min, \hat{\ell}}^{\mathcal{A}} \geq V_{\min, j}^{\mathcal{A}}, j = \ell, \dots, \hat{\ell} - 1$  which establishes the equivalence of the expressions in (7.18) and (7.11) whenever  $\hat{\ell} > \ell$  as well.

Next, to prove the result in (7.19) we note that the result is trivially true when  $D \leq 0$  or  $n \in \mathcal{R}_L^{\mathcal{A}}$ . Hence, we assume that  $D > 0$  and  $n \notin \mathcal{R}_L^{\mathcal{A}}$ . Clearly, in this case we must have that  $n \in \mathcal{R}_L^{\mathcal{A} \cup \underline{e}}$ . Then notice that upon including  $\underline{e}$ , the best weighted rate on only RB  $n \in \mathcal{N}_{\ell}$  is improved, whereas the best weighted rates on all RBs in all other subframes, i.e., RBs in the set  $\mathcal{N} \setminus \{n\}$ , are not changed. Consequently, it can be verified that the subframe  $\hat{\ell}$  will remain a bottleneck subframe and the selection of RBs on subframes  $\hat{\ell} + 1, \dots, L$  will not change, i.e.,  $\mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=\hat{\ell}+1}^L \mathcal{N}_j) = \mathcal{R}_L^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=\hat{\ell}+1}^L \mathcal{N}_j)$ . Therefore, we have to compare  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A}} = \mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j)$  and  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A} \cup \underline{e}} = \mathcal{R}_L^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j)$ , keeping in mind that their cardinalities are identical. Again invoking the fact that the best weighted rates on the first  $\ell - 1$  subframes are not changed, we can deduce that the number of RBs assigned to the first  $\ell - 1$  subframes in  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A} \cup \underline{e}}$  must be no greater than that in  $\mathcal{R}_{\hat{\ell}}^{\mathcal{A}}$ , which further implies that

$$\mathcal{R}_{\hat{\ell}}^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=1}^{\ell-1} \mathcal{N}_j) \subseteq \mathcal{R}_{\hat{\ell}}^{\mathcal{A}} \cap (\cup_{j=1}^{\ell-1} \mathcal{N}_j).$$

Furthermore, we can deduce that if

$$|\mathcal{R}_{\hat{\ell}}^{\mathcal{A}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j)| \leq |\mathcal{R}_{\hat{\ell}}^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j)|,$$

then

$$\mathcal{R}_\ell^{\mathcal{A}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j) \subseteq \mathcal{R}_\ell^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j),$$

and vice versa. Similarly at subframe  $\ell$  since we improve the best weighted rate at one RB  $n \in \mathcal{N}_\ell$ , we have that if  $|(\mathcal{R}_\ell^{\mathcal{A} \cup \underline{e}} \cap \mathcal{N}_\ell) \setminus \{n\}| \geq |\mathcal{R}_\ell^{\mathcal{A}} \cap \mathcal{N}_\ell|$  then  $\mathcal{R}_\ell^{\mathcal{A}} \cap \mathcal{N}_\ell \subseteq (\mathcal{R}_\ell^{\mathcal{A} \cup \underline{e}} \cap \mathcal{N}_\ell) \setminus \{n\}$  and vice versa. Next, invoking the weighted sum rate optimality of  $\mathcal{R}_\ell^{\mathcal{A}}$  (under the best weighted rates induced by the set of 3-tuples  $\mathcal{A}$ ), we can argue that

$$|\mathcal{R}_\ell^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=1}^{\ell-1} \mathcal{N}_j)| \geq |\mathcal{R}_\ell^{\mathcal{A}} \cap (\cup_{j=1}^{\ell-1} \mathcal{N}_j)| - 1.$$

This is because otherwise we can move an RB assignment from subframes  $\ell, \dots, \hat{\ell}$  to the first  $\ell - 1$  subframes without reducing (or with improving) the weighted sum rate while retaining feasibility, until the relation is satisfied. The same argument also allows us to conclude that  $|\mathcal{R}_\ell^{\mathcal{A} \cup \underline{e}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j)| \leq |\mathcal{R}_\ell^{\mathcal{A}} \cap (\cup_{j=\ell+1}^{\hat{\ell}} \mathcal{N}_j)| + 1$ . Together, these observations suffice to conclude that (7.19) holds true.  $\square$

Finally, we note that Algorithm 3 can be initialized with  $V_n^\emptyset = 0, \forall n \in \mathcal{N}$  and any arbitrary choice of  $\mathcal{R}_L^\emptyset$  satisfying the cardinality and causality constraints  $|\mathcal{R}_L^\emptyset \cap \{\cup_{j=1}^\ell \mathcal{N}_j\}| = \min\{J, \sum_{j=1}^\ell J_j\}, \forall \ell = 1, \dots, L$ .

---

### Algorithm 3

---

- 1: Initialize with set of modes,  $\mathcal{M}$ , selected set of 3-tuples,  $\hat{\mathcal{G}} = \phi$ , and user set  $\mathcal{U}' = \mathcal{U}$ .
  - 2: **Repeat**
  - 3: Invoke Algorithm 4 to determine  $\hat{\underline{e}}$  as the tuple in  $\Psi \setminus \hat{\mathcal{G}}$  which offers (approximately) the largest gain among all tuples  $\underline{e} \in \Psi \setminus \hat{\mathcal{G}}$  such that  $\hat{\mathcal{G}} \cup \underline{e} \in \underline{\mathcal{I}}$ .
  - 4: **If**  $G = h(\hat{\mathcal{G}} \cup \hat{\underline{e}}) - h(\hat{\mathcal{G}}) > 0$  then
  - 5: Update  $\hat{\mathcal{G}} = \hat{\mathcal{G}} \cup \hat{\underline{e}}$  and  $\mathcal{U}' = \mathcal{U}' \setminus \{\hat{u}\}$ , where  $\hat{\underline{e}} = (\hat{u}, \hat{m}, \hat{\mathbf{b}})$ .
  - 6: **End If**
  - 7: **Until**  $\mathcal{U}' = \phi$  or  $G = 0$ .
  - 8: Output  $\hat{\mathcal{G}}$ .
- 

## 7.3 Extensions

In this section we briefly comment on extensions of our results to a multi-cell setting where inter-cell interference coordination is important. Our first observation is that any given time-frequency resource partition among cells in which each cell is allowed to only use a subset of the available RBs in a scheduling block, can be readily accommodated.

---

**Algorithm 4**


---

- 1: Initialize with user set,  $\mathcal{U}'$ , and set of modes,  $\mathcal{M}$  and set of selected tuples  $\hat{\mathcal{G}}$
- 2: **For** each user  $u \in \mathcal{U}'$  and each mode  $m \in \mathcal{M}$  **do**
- 3: Define  $\hat{e} = (\hat{u}, \hat{m}, \hat{\mathbf{b}}_u^m)$  with  $\hat{u} = u, \hat{m} = m, \hat{\mathbf{b}}_u^m = \mathbf{0}$  and set  $w = 0$  and gain  $G = 0$ .
- 4: **While**  $w < Q_u$  **Do**
- 5: Determine  $\check{e} = (u, m, \check{\mathbf{b}}_u^m) \in \mathcal{E}$  as

$$\arg \max_{\substack{\mathbf{e}=(u,m,\mathbf{b}_u^m) \in \mathcal{E} \\ \sum_{n \in \mathcal{N}} b_{u,n}^m > 0; w + \sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u}} \left\{ \frac{h(\hat{\mathcal{G}} \cup \hat{e} \cup \mathbf{e}) - h(\hat{\mathcal{G}} \cup \hat{e})}{\sum_{n \in \mathcal{N}} b_{u,n}^m} \right\} \quad (7.21)$$

- 6: Update  $\hat{b}_{u,\check{n}}^m = \check{b}_{u,\check{n}}^m$  where  $\check{b}_{u,\check{n}}^m > 0$  at  $n = \check{n}$  and  $\check{b}_{u,n}^m = 0, \forall n = \mathcal{N} \setminus \check{n}$
- 7: Update  $G = G + h(\hat{\mathcal{G}} \cup \hat{e} \cup \check{e}) - h(\hat{\mathcal{G}} \cup \hat{e})$  and  $w = w + \check{b}_{u,\check{n}}^m$ .
- 8: **End While**
- 9: Determine

$$\check{e} = \arg \max_{\substack{\mathbf{e}=(u,m,\mathbf{b}_u^m) \in \mathcal{E} \\ \sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u}} \left\{ h(\hat{\mathcal{G}} \cup \mathbf{e}) - h(\hat{\mathcal{G}}) \right\} \quad (7.22)$$

- 10: **If**  $G < h(\hat{\mathcal{G}} \cup \check{e}) - h(\hat{\mathcal{G}})$
  - 11: Update  $\hat{e} = \check{e}$  and  $G = h(\hat{\mathcal{G}} \cup \check{e}) - h(\hat{\mathcal{G}})$
  - 12: **End If**
  - 13: **End For**
  - 14: Output the 3-tuple  $\hat{e}$  along with its computed gain  $G$ , where that gain is largest among all computed gains over users in  $\mathcal{U}'$  and modes in  $\mathcal{M}$ .
- 

In particular, in each cell we can simply set the throughput to be zero for all users, modes and bit loadings, on each RB that is prohibited for that cell. Our second observation is more involved. We first note that per-subframe cardinality bounds can be imposed in order to limit the radiated energy per subframe (and hence the interference imposed on users served by other cells). In particular, we will show that additional per-subframe cardinality bounds can be placed on the formulation in (7.3) which can then be approximately solved via Algorithm 3. Indeed, let  $C_\ell, \ell = 1, \dots, L$  denote  $L$  per-subframe cardinality bounds. Using the same observation as in Lemma I we can assume without loss of generality that  $J_\ell \leq C_\ell \leq N, \forall \ell$  and consider the following



extension of (7.3)

$$\max_{\{x_{u,n}^m, b_{u,n}^m\}} \sum_{u=1}^K \sum_{\ell=1}^L \sum_{n \in \mathcal{N}_\ell} \sum_{m=1}^M \psi'_{u,n} x_{u,n}^m b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho)) \quad (7.23a)$$

$$\text{subject to } \sum_{m=1}^M \sum_{u=1}^K x_{u,n}^m \leq 1, \quad n \in \mathcal{N}, \quad (7.23b)$$

$$\sum_{u=1}^K \sum_{m=1}^M \max_{n \in \mathcal{N}} x_{u,n}^m \leq \bar{K}, \quad (7.23c)$$

$$\sum_{m=1}^M \max_{n \in \mathcal{N}} x_{u,n}^m \leq 1, \quad 1 \leq u \leq K, \quad (7.23d)$$

$$\sum_{n \in \mathcal{N}_1 \cup \dots \cup \mathcal{N}_\ell} \sum_{u=1}^K \sum_{m=1}^M x_{u,n}^m \leq \min \left\{ \sum_{q=1}^{\ell} J_q, J \right\}, \forall \ell, \quad (7.23e)$$

$$\sum_{n \in \mathcal{N}_\ell} \sum_{u=1}^K \sum_{m=1}^M x_{u,n}^m \leq C_\ell, \forall \ell, \quad (7.23f)$$

$$\sum_{m=1}^M \sum_{n \in \mathcal{N}} b_{u,n}^m \leq Q_u, \quad 1 \leq u \leq K, \quad (7.23g)$$

$$x_{u,n}^m \in \{0, 1\} \text{ \& } b_{u,n}^m \in \mathcal{B}^m, \forall u, m, n. \quad (7.23h)$$

Defining the set of 3-tuples,  $\Psi$ , as before, we can reformulate (7.23) as in (7.8) but where

$$\begin{aligned} h(\mathcal{A}) = & \max_{\substack{z_n \in \{0,1\} \\ \forall n}} \left\{ \sum_{n \in \mathcal{N}} z_n \max_{(u,m,b_u^m) \in \mathcal{A}} \{ \psi'_{u,n} b_{u,n}^m (1 - p_{u,n}^m(b_{u,n}^m, \rho)) \} \right\} \\ \text{s.t. } & \sum_{n \in \cup_{j=1}^{\ell} \mathcal{N}_j} z_n \leq \min \left\{ J, \sum_{q=1}^{\ell} J_q \right\}, \\ & \sum_{n \in \mathcal{N}_\ell} z_n \leq C_\ell, \quad \forall \ell = 1, \dots, L. \end{aligned} \quad (7.24)$$

It can be shown that Theorem 7.1 applies in this case as well and Algorithm 3 when initialized with a feasible choice of  $\mathcal{R}_L^\emptyset$  that satisfies  $|\mathcal{R}_L^\emptyset \cap \{\cup_{j=1}^{\ell} \mathcal{N}_j\}| = \min\{J, \sum_{j=1}^{\ell} J_j\}$ ,  $|\mathcal{R}_L^\emptyset \cap \mathcal{N}_\ell| \leq C_\ell \forall \ell$ , yields exactly the same guarantee as in Theorem 7.2. Finally, in implementing Algorithm 3 we can use Proposition 7.3 after a simple change in (7.18)

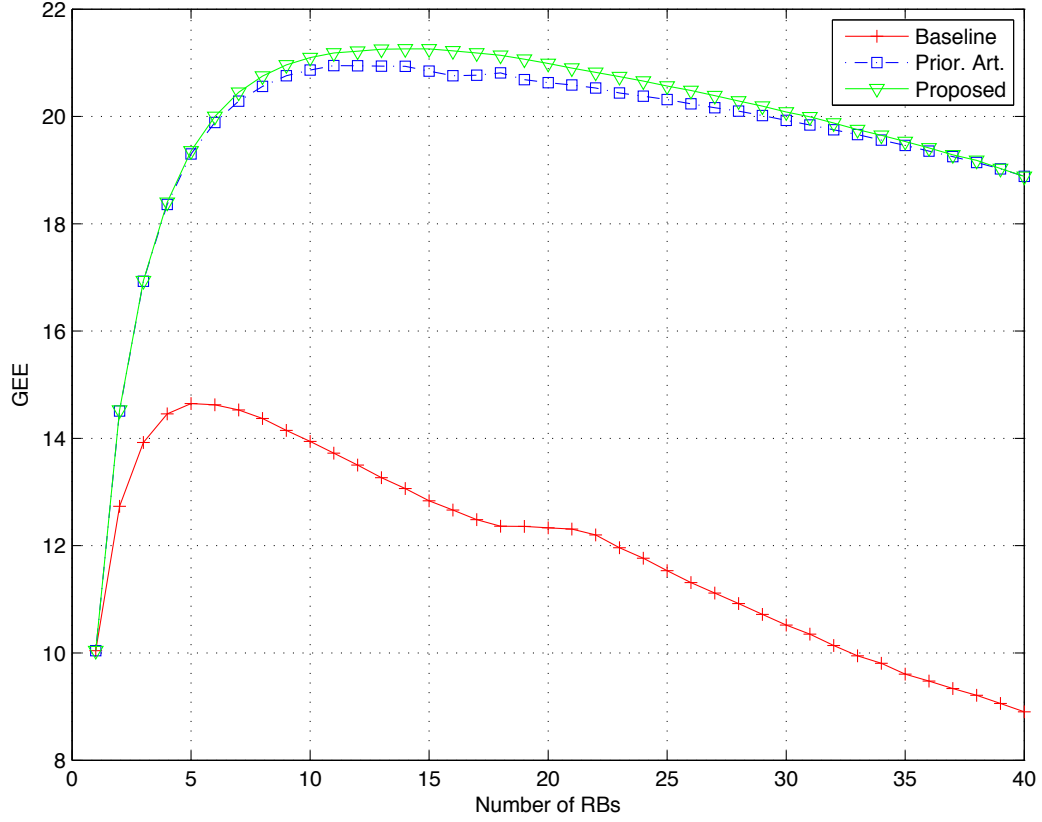


Figure 7.1: Achieved GEE versus Number of used RBs: Large Queue-sizes

to

$$V_n^{\mathcal{A}} = \max \left\{ \Delta_n^{\mathcal{A}}, \Delta_{\ell, C_\ell}^{\mathcal{A}}, \min_{s \in \mathcal{R}_L^{\mathcal{A}} \cap (\cup_{j=1}^{\hat{\ell}} \mathcal{N}_j)} \{ \Delta_s^{\mathcal{A}} \} \right\}, \quad (7.25)$$

where  $\Delta_{\ell, C_\ell}^{\mathcal{A}}$  is the  $(C_\ell)^{th}$  largest member of  $\{\Delta_n^{\mathcal{A}}\}_{n \in \mathcal{N}_\ell}$  with  $\hat{n}_\ell$  denoting the corresponding RB and  $\hat{\ell} = \min\{j \in \mathcal{S}^{\mathcal{A}} : j \geq \ell\}$ . Similarly, upon inclusion of any new 3-tuple  $\underline{e} = (u, m, \mathbf{b}_u^m) \in \mathcal{E}$ , such that  $b_{u,n}^m > 0$  for that  $n \in \mathcal{N}_\ell$ , the set  $\mathcal{R}_L^{\mathcal{A} \cup \underline{e}}$  can be determined using

$$\mathcal{R}_L^{\mathcal{A} \cup \underline{e}} = \begin{cases} \mathcal{R}_L^{\mathcal{A}} & D \leq 0 \text{ or } n \in \mathcal{R}_L^{\mathcal{A}}, \\ (\mathcal{R}_L^{\mathcal{A}} \setminus \{\hat{n}_\ell\}) \cup \{n\} & V_n^{\mathcal{A}} = \Delta_{\ell, C_\ell}^{\mathcal{A}}, \\ (\mathcal{R}_L^{\mathcal{A}} \setminus \{\hat{s}\}) \cup \{n\} & \text{else.} \end{cases} \quad (7.26)$$

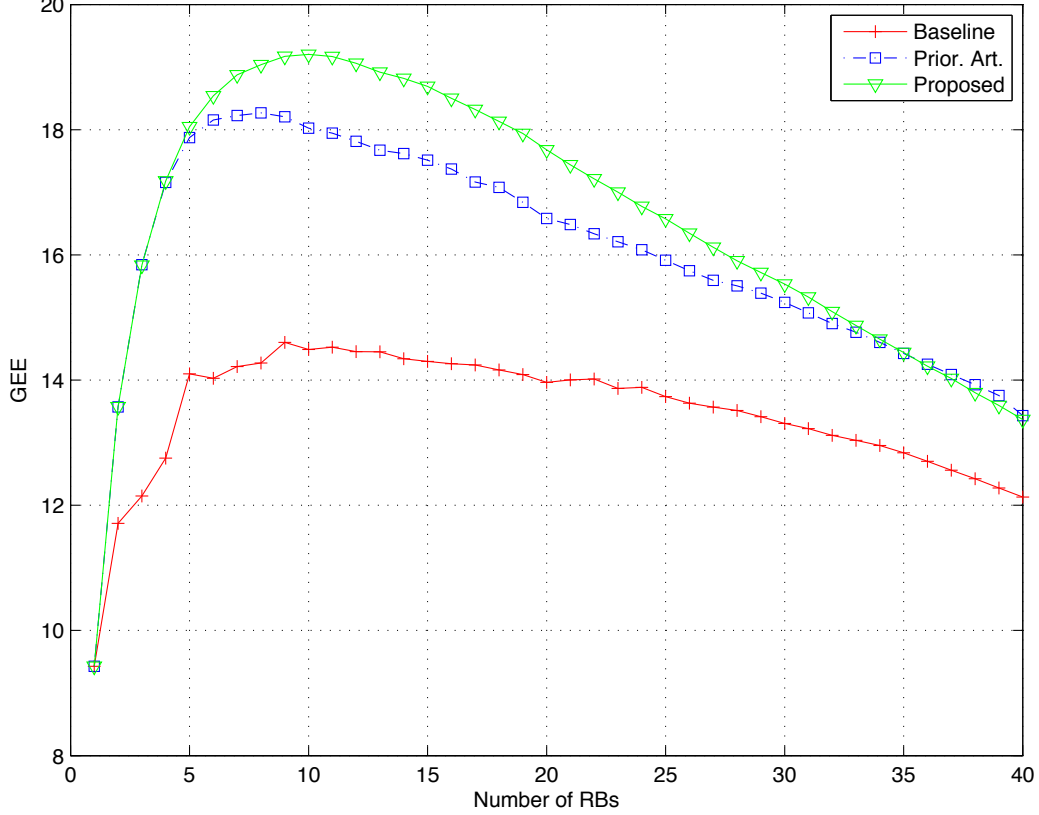


Figure 7.2: Achieved GEE versus Number of used RBs: Small Queue-sizes

## 7.4 Simulation Results

We conducted a numerical study by evaluating the achieved average GEE over 1000 random input traces and the 0 – 1 throughput model. For simplicity, we first consider a scheduling block length of one ( $L = 1$ ). This allows us to compare our algorithm with two benchmarks that entail deterministic algorithms of comparable complexity, and which have been suggested for constrained weighted sum rate maximization (subframe scheduling). We will later consider the more general multi-subframe per block case. In our study, we fixed all RB weights,  $\{\psi_n\}$ , to be identical, the number of modes,  $M$ , to be 4 and the number of RBs,  $N$ , to be 40. The number of users,  $K$ , as well as user limit  $\bar{K}$  were chosen to be 20. The usable RB set cardinality,  $J$ , was varied from 1 to  $N$ . We chose one representative value for the per-RB energy  $\rho$  along with values for the harvested energy per subframe  $E_1$  and the baseband circuit energy  $\vartheta$  (the power amplifier efficiency was chosen to be 1) which resulted in  $J_1 = N$ . We note that these

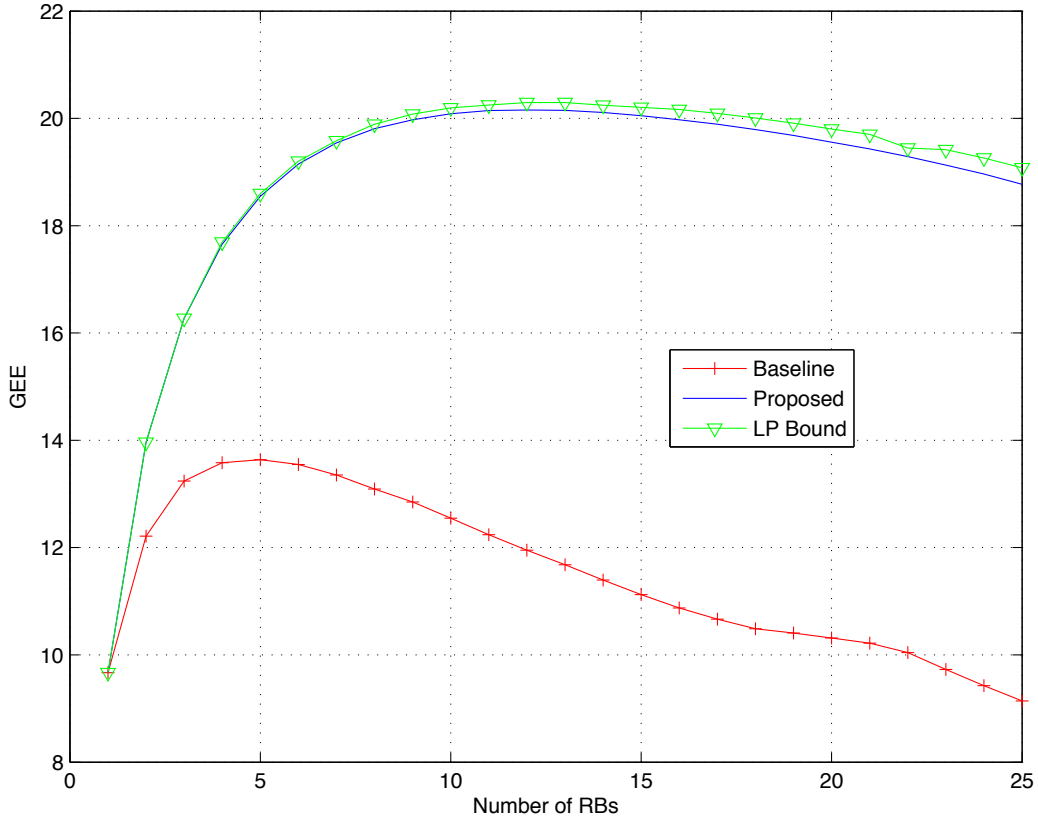


Figure 7.3: Achieved GEE versus Number of used RBs: Large Queue-sizes

chosen values were fixed across all the 1000 traces. On the other hand, the per-user per-RB rates were generated in an i.i.d. manner across traces. In the results presented here these rates were generated using the half-normal distribution. Similar trends were observed for other distributions.

The first benchmark we compare against is an extension of the greedy algorithm from [57, 59] that incorporates an additional knapsack constraint to limit the number of used RBs to some specified  $J$ . In particular, this extension follows the approach of [57, 59] (which we remind does not consider a limit on the number of usable RBs), until the RB cardinality constraint is reached. The other benchmark can accommodate any arbitrary linear cost constraint on the set of used RBs and proposes an approximation algorithm whose approximation guarantee decays logarithmically in the system dimension [87]. Compared to the algorithm in [87], the one proposed here offers a superior constant factor approximation guarantee. Moreover, while the algorithm proposed in [87] can

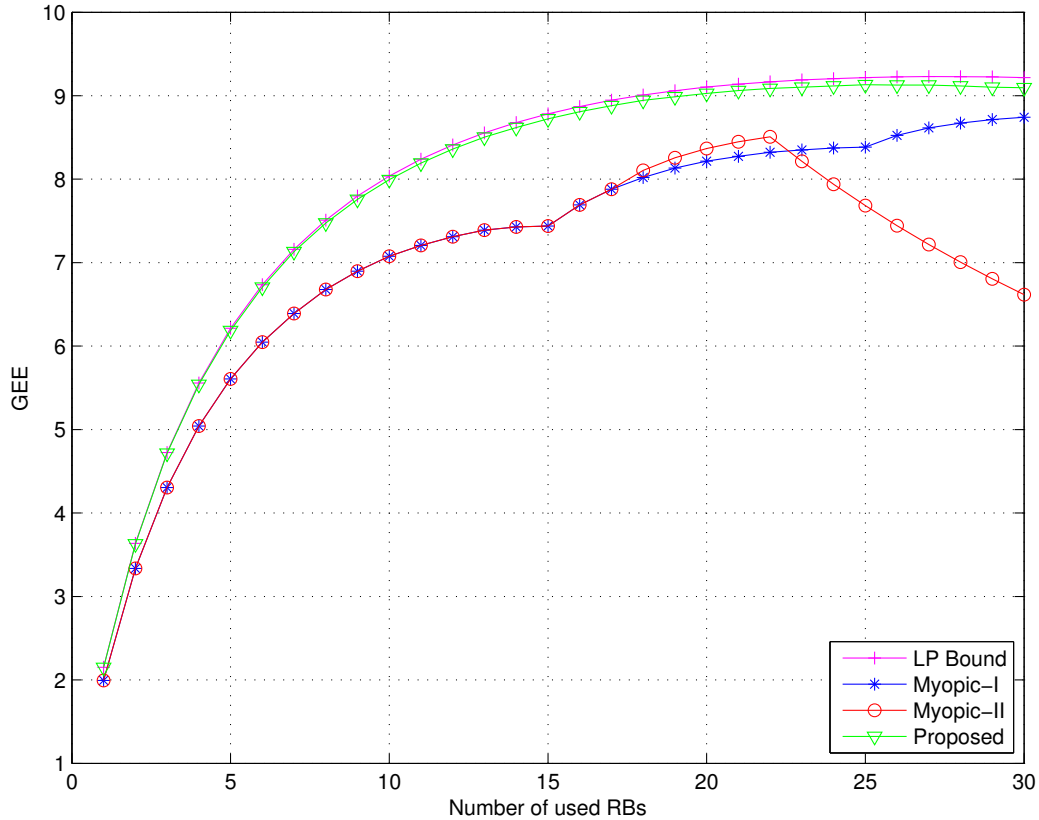


Figure 7.4: Achieved GEE versus Number of used RBs: Large Queue-sizes and  $L = 3$  subframes per block

incorporate a generic linear constraint, it cannot incorporate multiple linear causality constraints as the one proposed here.

In Fig. 7.1 we consider the large queue-size regime and plot the energy efficiency achieved (defined as the ratio of the weighted sum rate (computed per subframe) and the energy consumed per subframe) by our proposed algorithm, as well as the aforementioned two benchmarks, for each value of  $J = 1, \dots, N$ . Fig. 7.2 is the counterpart of Fig. 7.1 under a small queue size regime. From the figures it is seen that our proposed algorithm is significantly superior to the direct extension of the existing greedy methods (Baseline) and also outperforms the one from [87] (Prior art). Next, to assess the gap to optimality, we compare the performance offered by our algorithm with a linear programming (LP) based upper bound to (7.3) that can be obtained after some manipulations upon relaxing the binary indicator variables. In Fig. 7.3 we plot the performance (GEE) achieved by our algorithm and the upper bound (which we remind

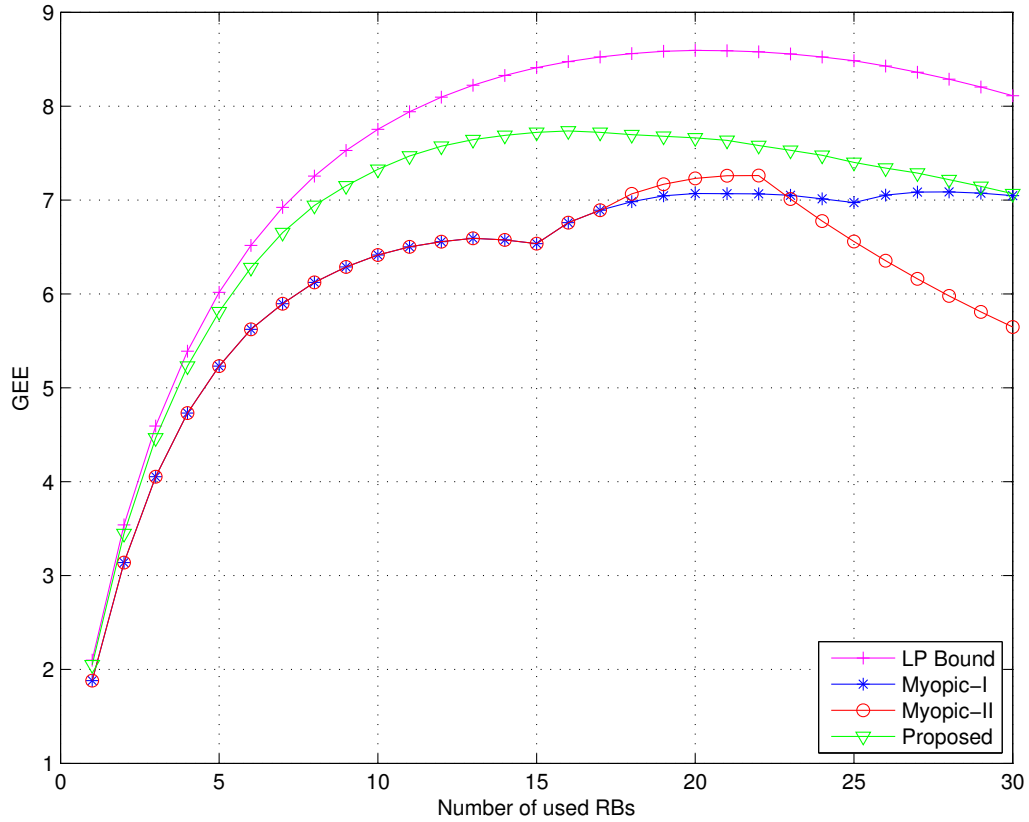


Figure 7.5: Achieved GEE versus Number of used RBs: Small Queue-sizes and  $L = 3$  subframes per block

need not be achievable). It is evident that our algorithm performs quite close to the optimal.

Finally, to benchmark the performance of our proposed algorithm over a more general case, we let each scheduling block have a length of  $L = 3$  subframes, with each subframe comprising of  $N = 15$  RBs. We assume that the circuit power is supplied by a non-renewable source and that causality constraints are specified via parameters  $J_1 = 15, J_2 = 10$  &  $J_3 = 5$ . The other parameters are chosen as in the previous examples. Then, in Fig. 7.4 and Fig. 7.5 we plot the performance (GEE) achieved by our algorithm and the LP upper bound for the large and small queue size regimes, respectively. Also plotted in these two figures are the performance of two myopic policies denoted by Myopic-I and Myopic-II, respectively. In Myopic-I we simulate a variation of the “*spend what you get*” policy [88]. At each subframe, this policy only exploits the CSI pertaining to that subframe and seeks to use the maximum number of RBs

possible under the energy available at the subframe. This policy is suitable for scenarios in which there is strict causality in the availability of CSI (and the evolution of CSI across subframes cannot be easily modeled). This policy utilizes the battery for storing energy that cannot be utilized, where such an excess energy remains when the available energy cannot be exhausted even upon using all the  $N$  RBs in a subframe. Recall from Lemma 7.1 that the parameters  $\{J_\ell\}_{\ell=1}^L$  we set are essentially the number of RBs this Myopic policy must use in each subframe, when not further constrained by a limit on the total number of usable RBs. To implement this policy, we first consider subframe 1 and specialize our proposed algorithm to a block length of one subframe with a cardinality bound  $\min\{J, J_1\}$  on the number of usable RBs. In each subsequent subframe,  $\ell$ , we obtain the scheduling decision by specializing Algorithm 3 to a block length of one subframe with a cardinality bound  $\min\{J - \sum_{j=1}^{\ell-1} J_j, J_\ell\}$  on the number of usable RBs and with the additional restrictions that:

1. Any user previously scheduled in any preceding subframe can only be scheduled in subframe  $\ell$  with the same previously assigned mode and cannot be assigned more than the remaining bits in its buffer.
2. The total number of distinct scheduled users up to subframe  $\ell$  cannot exceed  $\bar{K}$ .

We note that the latter two constraints are easy to incorporate. The other Myopic-II policy is the variant of the “*spend what you get*” policy which does not use the battery so that any excess energy needs to be discarded. For the Myopic-II policy we accordingly set the parameters as  $J_1 = 15, J_2 = 2, J_3 = 5$  to model the event that the energy arrival in subframe 2 is small and no excess energy from subframe 1 can be used in subframe 2. From these figures we can deduce that the performance of our algorithm is quite close to the optimal. Moreover, there is significant gain compared to the myopic policies especially when the total number of usable RBs  $J$  is limited. This is because since the myopic policies are committed to using as many RBs as possible in each subframe, they are unable to utilize many good RBs at the subsequent subframes. As expected both myopic policies have identical performance for  $J : 1 \leq J \leq 15$  since exactly the same scheduling decisions would be obtained for the first subframe under either policy.

Surprisingly Myopic-II can outperform Myopic-I for some values of  $J$  :  $16 \leq J \leq 22$ . Here, the fact that only two RBs can be used in the second subframe seems to be an advantage for the former policy because it forces the policy to consider the third subframe with fewer restrictions (recall that a user once scheduled can only be scheduled again under the same mode). For  $J > 22$ , Myopic-II policy degrades since it cannot exploit the battery to use more than 22 RBs. Notice that our proposed algorithm can be readily adapted to such myopic policies. Under the original non myopic policy, our algorithm judiciously considers all three subframes to obtain its scheduling decision and thus outperforms both myopic scheduling ones.

## 7.5 Conclusions

We proposed novel algorithms for optimizing global energy efficiency and weighted sum of energy efficiencies, respectively, over energy harvesting LTE OFDMA networks. The proposed algorithms were shown to guarantee constant-factor approximation and are simple enough to be implementable. An important avenue for future work is to extend our results for the multi-cell case.



## Chapter 8

### Conclusions and Future Work

In this thesis, we studied communication systems and protocols under energy constraints. At the receiver side, we investigated the fundamental limits of reliable communication when the processing is powered by random energy sources and subject to constraints on energy storage. This would be more vital in short-range communication applications where the receive energy is comparable to the transmit energy. We propose a model for the processing energy at the receiver that captures the trade-off between sampling energy and decoding energy. The model relies on the decoding energy being a decreasing function of the capacity gap between the code rate and the channel capacity. While sampling and decoding energies are typically comparable, the key issue is that the sampling is a real-time process; the samples must be collected during the transmission time of that packet. Thus the energy harvesting rate and battery size may limit the sampling rate. This model allows us to characterize the maximum throughput of a basic communication channel with limited processing energy.

We extended this result to fading channels and multi-user scenarios with limited processing energy at the receiver. We showed how to capture the receive multi-user diversity, in which the receiver decodes the messages experiencing the strongest channels in order to reduce the decoding energy per user and thus decode more data messages.

Next, we studied using HARQ protocol as a scheme to save energy at the receiver. Unlike conventional cases where the processing energy is not limited, here even if the receiver can decode a message, it may still ask the transmitter to send extra redundant bits in order to increase the capacity gap and reduce the processing energy. On the other hand, the extra redundant bits will increase the code length which in turn may increase the processing energy. Thus, the receiver requests retransmission as long as it reduces

the overall decoding energy. We also considered using Repetition-HARQ scheme along with MRC at the receiver. Although in this scheme the capacity is not improved as fast as in IR-HARQ, the key advantage is that the effective code length is kept constant. We show that in contrast to the traditional systems, here the Repetition-HARQ could outperform IR-HARQ.

We also investigated the problem of energy efficiency and energy harvesting in LTE networks. We formulated problems targeting maximization of weighted sum rate under either energy constraints or energy causality constraints imposed by energy harvesting devices. We showed the problems under consideration are NP-hard and proposed efficient algorithms to solve them approximately using submodular optimization. We also derived approximation guarantees for all scenarios.

Based on our work on energy-aware communication in both the transmitter and the receiver, we will seek to work in the following directions as the future work.

- *ARQ optimization:* As a future direction, we will seek to analyze the IR-HARQ under energy harvesting at the receiver to derive the optimal decision strategy including the optimal threshold on the energy. The challenge is that any decision on decoding a message or not decoding and requesting extra redundancy would affect the future decisions. This memory is introduced by the existence of the battery that can store the energy for the future use. The mathematical framework to deal with this problem involves a dynamic programming problem subject to several constraints which is not trivial to solve. The next step will be to consider the more efficient IR-HARQ scheme in which in each retransmission, a punctured code is transmitted. Here, in the first transmission, the systematic bits along with some redundant bits from a long code book are transmitted and in the next retransmissions, a subset of the redundant bits are selected randomly for transmission with a different power level.
- *Multi-user channels with limited processing energy:* To approach the capacity of multi-user channels, researchers have proposed a variety of signaling and coding techniques such as superposition coding, dirty-paper coding, successive and joint

decoding, etc. However, the computational complexity (or equivalently the processing energy) of these coding techniques, as a function of the parameters of the channel has not been well-understood. In addition, it is not clear if these techniques would enlarge the achievable rate region of the channel, if the processing energy at the transmitter or the receiver is limited. One research direction is to model and analyze the fundamental limits of communication over multi-user channels, with limited processing energy.

- *Low data rates with sporadic data:* In some applications, such as body sensors, the data rate is low. Using a code with a simple structure like repetition code would be good enough for such a system. In addition, as there is not always data for transmission, it is sent in irregular and scattered time intervals. So, the system stays idle for a long portion of the time. However, the receiver needs to spend some energy sensing the channel regularly to check for the new data. Designing efficient communication systems under these constraints is of practical interest.
- *Resource scheduling with energy harvesting for multi-cell networks:* The extensions made to enforce per-subframe cardinality bounds in Chapter 7 can be imposed in order to limit the radiated energy per subframe and hence the interference imposed on users served by other cells. Based on this fact, we could extend scheduling results under energy harvesting constraints to multi-cell LTE networks.
- *Resource scheduling based on receivers' processing energy (LTE):* In this thesis, we have investigated the resource allocation, including resource blocks (time-frequency blocks) assignment and mode (e.g., pre-coder matrix ) selection, in wireless cellular systems, when the power at the transmitter is limited and random. The same problem can be considered when the processing power at the receiver side is limited. In that case, we may assign resource blocks and data rates to maximize the throughput, when the processing energy is the bottleneck.

## References

- [1] J. Yang and S. Ulukus, "Optimal packet scheduling in an energy harvesting communication system," *IEEE Trans. Commun.*, vol. 60, no. 1, pp. 220–230, Jan. 2012.
- [2] K. Tutuncuoglu and A. Yener, "Optimum transmission policies for battery limited energy harvesting nodes," *IEEE Trans. Wireless Commun.*, vol. 11, no. 3, pp. 1180–1189, Mar. 2012.
- [3] O. Ozel, K. Tutuncuoglu, J. Yang, S. Ulukus, and A. Yener, "Transmission with energy harvesting nodes in fading wireless channels: Optimal policies," *IEEE J. Sel. Areas Commun.*, vol. 29, no. 8, pp. 1732–1743, Sep. 2011.
- [4] O. Ozel and S. Ulukus, "Achieving AWGN capacity under stochastic energy harvesting," *IEEE Trans. Info. Theory*, vol. 58, no. 10, pp. 6471–6483, 2012.
- [5] M. Khuzani, H. Saffar, E. Alian, and P. Mitran, "On optimal online power policies for energy harvesting with finite-state Markov channels," in *Proc. IEEE ISIT*, July 2013, pp. 1586–1590.
- [6] K. Tutuncuoglu, O. Ozel, A. Yener, and S. Ulukus, "Binary energy harvesting channel with finite energy storage," in *Proc. IEEE ISIT*, July 2013, pp. 1591–1595.
- [7] J. Yang, O. Ozel, and S. Ulukus, "Broadcasting with an energy harvesting rechargeable transmitter," *IEEE Trans. Wireless Commun.*, vol. 11, no. 2, pp. 571–583, Feb. 2012.
- [8] J. Yang and S. Ulukus, "Optimal packet scheduling in a multiple access channel with rechargeable nodes," in *Proc. IEEE ICC*, June 2011, pp. 1–5.
- [9] K. Tutuncuoglu and A. Yener, "Sum-rate optimal power policies for energy harvesting transmitters in an interference channel," *CoRR*, vol. abs/1110.6161, 2011.
- [10] O. Orhan and E. Erkip, "Throughput maximization for energy harvesting two-hop networks," in *Proc. IEEE ISIT*, July 2013, pp. 1596–1600.
- [11] P. Blasco, D. Gunduz, and M. Dohler, "Low-complexity scheduling policies for energy harvesting communication networks," in *Proc. IEEE ISIT*, July 2013, pp. 1601–1605.
- [12] C. Huang, R. Zhang, and S. Cui, "Throughput maximization for the gaussian relay channel with energy harvesting constraints," *IEEE J. Sel. Areas Commun.*, vol. 31, no. 8, pp. 1469–1479, August 2013.

- [13] D. Gunduz and B. Devillers, “Two-hop communication with energy harvesting,” in *4th IEEE International Workshop on Computational Advances in Multi-Sensor Adaptive Processing (CAMSAP)*, Dec 2011, pp. 201–204.
- [14] D. Gunduz, K. Stamatiou, N. Michelusi, and M. Zorzi, “Designing intelligent energy harvesting communication systems,” *IEEE Communications Magazine*, vol. 52, no. 1, pp. 210–216, January 2014.
- [15] O. Orhan, D. Gunduz, and E. Erkip, “Throughput maximization for an energy harvesting communication system with processing cost,” in *Information Theory Workshop (ITW), 2012 IEEE*, Sept 2012, pp. 84–88.
- [16] J. Xu and R. Zhang, “Throughput optimal policies for energy harvesting wireless transmitters with non-ideal circuit power,” *IEEE J. Sel. Areas Commun.*, vol. 32, no. 2, pp. 322–332, February 2014.
- [17] K. Tutuncuoglu and A. Yener, “Communicating with energy harvesting transmitters and receivers,” in *Information Theory and Applications Workshop (ITA), 2012*, Feb. 2012, pp. 240–245.
- [18] S. Cui, A. Goldsmith, and A. Bahai, “Energy-constrained modulation optimization,” *IEEE Trans. Wireless Commun.*, vol. 4, no. 5, pp. 2349–2360, Sept. 2005.
- [19] E. Lauwers and G. Gielen, “Power estimation methods for analog circuits for architectural exploration of integrated systems,” *IEEE Trans. Very Large Scale Integration (VLSI) Systems*, vol. 10, no. 2, pp. 155–162, Apr. 2002.
- [20] H. Mahdavi-Doost and R. Yates, “Energy harvesting receivers: Finite battery capacity,” in *Proc. IEEE ISIT*, 2013.
- [21] —, “Fading channels in energy-harvesting receivers,” in *Proc. CISS*, 2014.
- [22] R. Yates and H. Mahdavi-Doost, “Energy harvesting receivers: Optimal sampling and decoding policies,” in *IEEE Global Conference on Signal and Information Processing (GlobalSIP)*, Dec 2013, pp. 367–370.
- [23] H. Mahdavi-Doost and R. Yates, “Opportunistic reception in a multiuser slow-fading channel with an energy harvesting receiver,” in *Proc. IEEE ICC*, 2015.
- [24] H. Mahdavi-Doost and R. D. Yates, “Hybrid ARQ in block-fading channels with an energy harvesting receiver,” in *Proc. IEEE ISIT*, June 2015, pp. 1144–1148.
- [25] O. Ozel, K. Tutuncuoglu, S. Ulukus, and A. Yener, “Capacity of the energy harvesting channel with energy arrival information at the receiver,” in *Information Theory Workshop (ITW), 2014 IEEE*, Nov 2014, pp. 331–335.
- [26] A. Arafa and S. Ulukus, “Optimal policies for wireless networks with energy harvesting transmitters and receivers: Effects of decoding costs,” *IEEE J. Sel. Areas Commun.*, vol. 33, no. 12, pp. 2611–2625, Dec 2015.
- [27] A. Khandekar and R. McEliece, “On the complexity of reliable communication on the erasure channel,” in *Proc. IEEE ISIT*, 2001.

- [28] M. Luby, M. Mitzenmacher, M. Shokrollahi, and D. Spielman, "Efficient erasure correcting codes," *IEEE Trans. Info. Theory*, vol. 47, no. 2, pp. 569–584, feb 2001.
- [29] T. Richardson, M. Shokrollahi, and R. Urbanke, "Design of capacity-approaching irregular low-density parity-check codes," *IEEE Trans. Info. Theory*, vol. 47, no. 2, pp. 619–637, Feb 2001.
- [30] I. Sason and R. Urbanke, "Complexity versus performance of capacity-achieving irregular repeat-accumulate codes on the binary erasure channel," *IEEE Trans. Info. Theory*, vol. 50, no. 6, pp. 1247–1256, June 2004.
- [31] T. Richardson and R. Urbanke, "The renaissance of Gallager's low-density parity-check codes," *IEEE Communications Magazine*, vol. 41, no. 8, pp. 126–131, Aug. 2003.
- [32] P. Grover, K. Woyach, and A. Sahai, "Towards a communication-theoretic understanding of system-level power consumption," *IEEE J. Sel. Areas Commun.*, vol. 29, no. 8, pp. 1744–1755, Sept. 2011.
- [33] P. Grover, A. Goldsmith, and A. Sahai, "Fundamental limits on the power consumption of encoding and decoding," in *Proc. IEEE ISIT*, July 2012, pp. 2716–2720.
- [34] S. Verdú and T. Weissman, "The information lost in erasures," *IEEE Trans. Info. Theory*, vol. 54, no. 11, pp. 5030–5058, nov. 2008.
- [35] D. Forney, "Concatenated codes," Ph.D. dissertation, M.I.T., 1965.
- [36] I. Sason, "On universal properties of capacity-approaching ldpc code ensembles," *IEEE Trans. Info. Theory*, vol. 55, no. 7, pp. 2956–2990, 2009.
- [37] V. Guruswami and P. Xia, "Polar codes: Speed of polarization and polynomial gap to capacity," in *IEEE Proc. FOCS*, 2013, pp. 310–319.
- [38] N. Alon and M. Luby, "A linear time erasure-resilient code with nearly optimal recovery," *IEEE Trans. Info. Theory*, vol. 42, no. 6, pp. 1732–1736, Nov 1996.
- [39] MATLAB, "Low-density parity-check codes from DVB-S.2 standar," <http://www.mathworks.com/help/comm/ref/dvbs2ldpc.html>.
- [40] S. Information Theory at the Chalmers University of Technology, Gothenburg, "Generation of LDPC codes," [http://itpp.sourceforge.net/4.3.1/ldpc\\_gen\\_codes.html](http://itpp.sourceforge.net/4.3.1/ldpc_gen_codes.html), 2013.
- [41] M. Mansour and N. Shanbhag, "A 640-mb/s 2048-bit programmable ldpc decoder chip," *IEEE Journal of Solid-State Circuits*, vol. 41, no. 3, pp. 684–698, March 2006.
- [42] M. K. Roberts and R. Jayabalan, "An area efficient and high throughput multi-rate quasi-cyclic LDPC decoder for IEEE 802.11n applications," *Microelectronics Journal*, no. 0, pp. –, 2014. [Online]. Available: <http://www.sciencedirect.com/science/article/pii/S0026269214002225>

- [43] W. Hoeffding, "Probability inequalities for sums of bounded random variables," *Journal of the American Statistical Association*, vol. 58, no. 301, pp. 13–30, March 1963. [Online]. Available: <http://www.jstor.org/stable/2282952?>
- [44] L. Kleinrock, *Queueing Systems*. John Wiley and Sons, Inc., 1975.
- [45] R. C. Larson and A. R. Odoni, *Urban Operations Research*. Prentice-Hall, NJ, 1981.
- [46] R. G. Gallager, *Stochastic Processes: Theory for Applications*. Cambridge University Press, 2014.
- [47] D. Tse and P. Viswanath, *Fundamentals of Wireless Communication*. Cambridge University Press, 2005.
- [48] R. Knopp and P. Humblet, "Information capacity and power control in single-cell multiuser communications," in *Proc. IEEE ICC*, vol. 1, 1995, pp. 331–335 vol.1.
- [49] D. Tse and S. Hanly, "Multiaccess fading channels. I. polymatroid structure, optimal resource allocation and throughput capacities," *IEEE Trans. Info. Theory*, vol. 44, no. 7, pp. 2796–2815, 1998.
- [50] L. Li and A. Goldsmith, "Capacity and optimal resource allocation for fading broadcast channels .i. ergodic capacity," *IEEE Trans. Info. Theory*, vol. 47, no. 3, pp. 1083–1102, 2001.
- [51] J. M. Wozencraft and M. Horstein, "Coding for two-way channels," *Tech. Rep. 383, Res. Lab. Electron., MIT, Cambridge, MA*, Jan. 1961.
- [52] G. Caire and D. Tuninetti, "The throughput of hybrid-ARQ protocols for the gaussian collision channel," *IEEE Trans. Info. Theory*, vol. 47, no. 5, pp. 1971–1988, Jul 2001.
- [53] D. Chase, "Code combining—a maximum-likelihood decoding approach for combining an arbitrary number of noisy packets," *IEEE Trans. Commun.*, vol. 33, no. 5, pp. 385–393, May 1985.
- [54] D. Rowitch and L. Milstein, "On the performance of hybrid FEC/ARQ systems using rate compatible punctured turbo (RCPT) codes," *IEEE Trans. Commun.*, vol. 48, no. 6, pp. 948–959, June 2000.
- [55] E. Soljanin, R. Liu, and P. Spasojevic, "Hybrid arq with random transmission assignments," *DIMACS Series in Discrete Mathematics and Theoretical Computer Science*, 2003.
- [56] E. Visotsky, V. Tripathi, and M. Honig, "Optimum ARQ design: a dynamic programming approach," in *Proc. IEEE ISIT*, June 2003.
- [57] M. Andrews and L. Zhang, "Scheduling algorithms for multi-carrier wireless data systems," in *ACM Mobicom 2007*, Sep. 2007.
- [58] S. Lee, S. Choudhury, A. Khoshnevis, S. Xu, and S. Lu, "Downlink MIMO with frequency-domain packet scheduling for 3gpp lte," in *IEEE Infocom*, Brazil, 2009.

- [59] H. Zhang, N. Prasad, and S. Rangarajan, "MIMO downlink scheduling in LTE and LTE-advanced systems," in *Proc. 2012 IEEE INFOCOM Miniconference*, Orlando, FL, Mar. 2012.
- [60] H. Ahmed, K. Jagannathan, and S. Bhashyam, "Queue-aware optimal resource allocation for the LTE downlink with best M sub-band feedback," *IEEE Trans. Wireless Comm.*, To appear.
- [61] S. N. Donthi and N. B. Mehta, "Joint performance analysis of channel quality indicator feedback schemes and frequency-domain scheduling for LTE," *IEEE Trans. Vehicular Tech.*, Sept. 2011.
- [62] A. Zappone and E. Jorswieck, "Energy efficiency in wireless networks via fractional programming theory," *Foundations and Trends in Comm. and Info. Theory*, 2015.
- [63] A. Kwasinski and A. Kwasinski, "Integrating cross-layer LTE resources and energy management for increased powering of base stations from renewable energy," in *Proc. 2015 IEEE WiOpt*, 2015.
- [64] Z. Wang, V. Aggarwal, and X. Wang, "Optimal energy-bandwidth allocation for energy harvesting interference networks," in *Proc. 2014 IEEE ISIT*, 2014.
- [65] E. V. Belmega and S. Lasaulce, "Energy-efficient precoding for multiple antenna terminals," *IEEE Trans. on Signal Proc.*, Jan. 2011.
- [66] G. Miao, N. Himayat, and G. Y. Li, "Energy-efficient link adaptation in frequency-selective channels," *IEEE Trans. on Commun.*, Feb. 2010.
- [67] D. Ng, E. Lo, and R. Schober, "Energy-efficient resource allocation in ofdma systems with large numbers of base station antennas," *IEEE Trans. on Wireless Comm.*, Sep. 2012.
- [68] L. Venturino, A. Zappone, C. Risi, and S. Buzzi, "Energy-efficient scheduling and power allocation in downlink OFDMA networks with base station coordination," *IEEE Trans. on Wireless Comm.*, vol. 14, no. 1, pp. 1–14, Jan. 2015.
- [69] D. Ng, E. Lo, and R. Schober, "Energy-efficient resource allocation in multi-cell OFDMA systems with limited backhaul capacity," *IEEE Trans. on Commun.*, Oct. 2012.
- [70] L. Venturino and S. Buzzi, "Energy-aware and rate-aware heuristic beamforming in downlink MIMO OFDMA networks with base station coordination," *IEEE Trans. on Vehicular Tech.*, To appear 2015.
- [71] G. L. Nemhauser, L. A. Wolsey, and M. L. Fisher, "An analysis of approximations for maximizing submodular set functions-I," *Mathematical Programming*, July 1978.
- [72] Y. Azar and I. Gamzu, "Efficient submodular function maximization under linear packing constraints," in *39th Intl. Colloquium on Automata, Languages and Programming*, 2012.



- [73] R. Iyer and J. A. Bilmes, “Submodular optimization with submodular cover and submodular knapsack constraints,” *NIPS*, 2013.
- [74] D. Pisinger and P. Toth, “Knapsack problems,” *Handbook of Combinatorial Optimization*, 1998.
- [75] H. Lin and J. Bilmes, “Multi-document summarization via budgeted maximization of submodular functions,” in *HLT10: Human Language Technologies*, 2010, pp. 912–920.
- [76] H. Mahdavi-Doost, N. Prasad, and S. Rangarajan, “Energy efficient downlink scheduling in LTE-advanced networks,” Rutgers University and NEC Labs America, Tech. Rep. <https://www.dropbox.com/s/kblg09zjh9ylp70/EESLong.pdf?dl=0>, 2015.
- [77] U. Feige, “On maximizing welfare when utility functions are subadditive,” *STOC*, 2006.
- [78] X. Kang, Y.-K. Chia, C. K. Ho, and S. Sun, “Cost minimization for fading channels with energy harvesting and conventional energy,” *IEEE Trans. Wireless Commun.*, vol. 13, no. 8, pp. 4586–4598, Aug 2014.
- [79] D. Ng, E. Lo, and R. Schober, “Energy-efficient resource allocation in OFDMA systems with hybrid energy harvesting base station,” *IEEE Trans. Wireless Commun.*, vol. 12, no. 7, pp. 3412–3427, July 2013.
- [80] Z. Wang, X. Wang, and V. Aggarwal, “Energy-subchannel allocation for energy harvesting nodes in frequency-selective channels,” in *Proc. IEEE ISIT*, June 2015, pp. 2712–2716.
- [81] W. Zeng, Y. Zheng, and R. Schober, “Online resource allocation for energy harvesting downlink multiuser systems: Precoding with modulation, coding rate, and subchannel selection,” *IEEE Trans. Wireless Commun.*, vol. 14, no. 10, pp. 5780–5794, Oct 2015.
- [82] N. Tekbiyik, T. Girici, E. Uysal-Biyikoglu, and K. Leblebicioglu, “Proportional fair resource allocation on an energy harvesting downlink,” *IEEE Trans. Wireless Commun.*, vol. 12, no. 4, pp. 1699–1711, April 2013.
- [83] H. Li, J. Xu, R. Zhang, and S. Cui, “A general utility optimization framework for energy harvesting based wireless communications,” *CoRR*, vol. abs/1501.01345, 2015. [Online]. Available: <http://arxiv.org/abs/1501.01345>
- [84] 3GPP, “Evolved universal terrestrial radio access (E-UTRA)-radio resource control (RRC),” *TR36.331 V10.7.0*, Mar. 2013.
- [85] W. Yu and J. Cioffi, “Constant power water-filling: Performance bound and low-complexity implementation,” *IEEE Transactions on Communications*, vol. 54, no. 1, Jan. 2006.
- [86] P. Goundan and A. Schulz, “Revisiting the greedy approach to submodular set function maximization,” *manuscript*, Jun. 2007.

- [87] H. Mahdavi-Doost, N. Prasad, and S. Rangarajan, “Energy efficient downlink scheduling in lte-advanced networks,” in *8th International Conference on Communication Systems and Networks (COMSNETS)*, Jan 2016, pp. 1–8.
- [88] M. Gorlatova, A. Bernstein, and G. Zussman, “Performance evaluation of resource allocation policies for energy harvesting devices,” in *International Symposium on Modeling and Optimization in Mobile, Ad Hoc and Wireless Networks (WiOpt)*, May 2011, pp. 189–196.