

UNIVERSITY OF CALIFORNIA

Los Angeles

Scheduling Algorithms for Millimeter-Wave Full-Duplex Integrated Access and Backhaul

A thesis submitted in partial satisfaction

of the requirements for the degree

Master of Science in Electrical and Computer Engineering

by

Sungho Cho

2026

© Copyright by

Sungho Cho

2026

ABSTRACT OF THE THESIS

Scheduling Algorithms for Millimeter-Wave Full-Duplex Integrated Access and Backhaul

by

Sungho Cho

Master of Science in Electrical and Computer Engineering

University of California, Los Angeles, 2026

Professor Ian P. Roberts, Chair

Millimeter-wave integrated access and backhaul (IAB) enables dense wireless deployments by allowing access and backhaul links to share spectrum. Full-duplex operation can further improve spectral efficiency and reduce delay, but self-interference makes its practical scheduling problem challenging. In full-duplex IAB networks, each node must coordinate multiple antenna panels whose simultaneous transmissions affect receivers, making link rates schedule-dependent.

This thesis studies scheduling for full-duplex mmWave IAB networks under self-interference. We formulate the scheduling problem as a maximum-weight optimization over a multi-hop IAB tree and develop a forward-backward dynamic programming algorithm that computes the throughput-optimal schedule with nearly linear complexity. We also propose FD-LocalMW-BP, a local max-weight scheduler only based on local information. FD-LocalMW-BP learns effective backhaul-downlink weights during an estimation phase and reuses them during deployment to approximate unobservable self-interference-dependent rates.

We prove stability of the centralized back-pressure policy and derive a stability condition for FD-LocalMW-BP. Simulations show that FD-LocalMW-BP approaches global scheduling performance while substantially reducing control overhead, especially in self-interference-limited regimes.

The thesis of Sungho Cho is approved.

Danijela Cabric

Lieven Vandenberghe

Ian P. Roberts, Committee Chair

University of California, Los Angeles

2026

For my beloved family and friends.

Contents

Abstract	ii
List of Figures	vi
Acknowledgements	1
Acronyms	2
1 Introduction	3
1.1 Background and Motivation	3
1.2 Contributions	6
1.3 Organization	6
2 Background	7
2.1 Full-Duplex Operation and Self-Interference	7
2.2 Preliminaries	9
2.3 Notation and Conventions	11
3 Full-duplex Integrated Access Backhaul Scheduling	15
3.1 System Model	15
3.2 Dynamic Programming approach	25
3.3 Local Max-Weight Scheduling	33
4 Numerical Results	63
4.1 Simulation Setup	63
5 Conclusion	73

List of Figures

3.1	Timing diagram of global versus local scheduling on a tree of depth D . Global scheduling requires a forward pass and a backward pass before TX/RX can begin, consuming $2D$ control intervals per slot. Local scheduling propagates only top-down consuming D control intervals per slot.	38
3.2	Operational timeline of FD-LocalMW-BP. After UE acquisition, the network runs the global policy at the inflated load $\lambda = \eta +$ virtual traffic for T_{est} slots (estimation phase) to produce the BH-DL weight estimates $\{\hat{\mu}_{b_e}\}$. The network then switches to FD-LocalMW-BP at the real load η (deployment phase), using $\hat{\mu}$ as the fixed BH-DL weight.	39
4.1	Simulated IAB topology.	65
4.2	Mean delay versus arrival rate λ , at three self-interference levels.	66
4.3	Mean delay versus self-interference (INR), at three arrival rates.	67
4.4	Mean UE delay versus arrival rate λ at INR = 0 dB, shown separately for each tree depth.	69
4.5	CDF of individual delay at $\lambda = 120$ Mbps and INR = 0 dB.	70
4.6	Distribution of the difference between the estimated weight $\hat{\mu}_{b_n}$ and the backhaul-DL rate $\bar{R}_q^{(n)}$ over all active backhaul-DL.	71
4.7	Approximated $\max_n \zeta_n$ of 3.122.	72

Acknowledgements

This thesis is made possible by the support and teachings from many people I had the honor of studying and interacting with during my time as a student at UCLA. Firstly, I would like to extend sincerest thanks to Professor Ian Roberts for his advisorship. His philosophy toward research and attitudes has been astoundingly insightful and fantastic, for which I will keep following that. I would also like to thank my committee members, Professor Lieven Vandenberghe, whose courses have been a big inspiration for my continued studies in optimization. I am also grateful to Professor Danijela Cabric for giving a valuable review of the thesis. Learning from these wonderful scholars will be the unforgettable experience and act as a continued motivation for my research.

Moreover, I am thankful for all my friends who have been great support, energy, and motivation. I cannot forget the moments with them, and without them, I was not able to overcome all hardships I have encountered in my life.

Finally, I am appreciative of my family, who have enthusiastically supported me through my educational journey and have been a continuous source of encouragement.

List of Acronyms

5G fifth-generation

6G sixth-generation

3GPP 3rd Generation Partnership Project

ADC analog-to-digital converter

dB decibel

FD full-duplex

HD half-duplex

SI Self-interference

SIC Self-interference cancellation

Chapter 1

Introduction

1.1 Background and Motivation

The growth of mobile data traffic, driven by immersive media, vehicular communication, and machine-to-machine connectivity, has placed steady pressure on cellular networks to deliver both higher capacity and broader coverage [2]. Millimeter-wave frequencies have become a cornerstone of fifth-generation and beyond systems because of their abundant spectrum and their support for highly directional transmission. The narrow beams of dense millimeter-wave (mmWave) arrays isolate users in space, suppressing user-to-user interference and enabling the aggressive spectral reuse that was not available at sub-6 GHz frequencies [13, 20].

Realizing mmWave at scale is held back by the cost of dense fiber backhaul. Integrated Access and Backhaul (IAB) addresses this by reusing the same mmWave spectrum for both user access and inter-node backhaul, so an operator can extend coverage by adding self-backhauled relay nodes instead of laying fiber. 3GPP standardized IAB in Releases 16 and 17 [1, 12], and the resulting multi-hop topologies, in which infrastructure nodes act at once as base stations and as wireless relays, are now a baseline assumption in mmWave network design.

The multi-hop structure that makes IAB attractive also creates its central difficulty.

Traffic for a deep node is relayed hop by hop along the backhaul, so a node near the donor must carry not only its own access traffic but the aggregated traffic of its entire subtree. As the tree grows, the backhaul links closer to the donor see proportionally more packets, and the load they must clear scales with the size of the subtree they feed. Whether the network remains stable therefore depends on how efficiently each node divides its radio resources between serving its own users and relaying traffic deeper into the tree.

A promising way to raise this efficiency is full-duplex (FD) operation, in which a node transmits and receives at the same time and frequency. An IAB node that can receive backhaul on one sector while transmitting access on another no longer wastes a slot waiting for the opposite direction, which directly relieves the relaying bottleneck above. Recent progress in self-interference cancellation, combining antenna isolation, analog and digital cancellation, and the spatial isolation of mmWave beamforming, has made FD operation practical enough to consider at the network layer [15, 23].

Two structural features of real mmWave base stations are often left out of scheduling models but turn out to be decisive here. First, these base stations are sectorized, hosting several antenna panels that point in different directions, so a node can run several links at once rather than a single transmit or receive beam [13]. Second, the cancellation of self-interference is never perfect, so when one panel transmits while another receives, a residual self-interference couples the two and lowers the rate of the receiving link by an amount that grows with the number of co-active transmit panels [14]. A scheduler that ignores either the multiple panels or the residual coupling between them does not describe the system it is meant to control.

The open problem, then, is link scheduling for a sectorized, full-duplex, multi-hop IAB network under residual self-interference. In each slot the scheduler must decide, across all panels of every node, which links transmit and which receive, with the rate of each receive link depending on how many panels at the same node transmit, and with decisions at one node propagating as load to the next.

1.1.1 Literature Reviews

Prior scheduling work for IAB and related mmWave networks leaves this problem open along five axes.

First, most models treat each node as carrying a single transmit or receive beam and do not represent the multiple panels of a sectorized base station, so they cannot express the choice of how many panels to commit to backhaul against access in a given slot.

Second, the studies that do consider full-duplex operation largely assume self-interference away. [22] designs separate FD and HD schedulers under ideal cancellation, [21] admits FD only under a no-interference model, and [5] proves throughput optimality for heterogeneous HD and FD networks but again assumes perfect cancellation and a single-hop topology. Because full-duplex is treated as interference-free, none of these can capture how residual self-interference shapes the decision.

Third, much of the literature retains the half-duplex assumption altogether [3, 9, 17, 24], in which a node transmits or receives but never both, and so forgoes the concurrency that motivates a sectorized full-duplex design in the first place.

Fourth, the throughput-optimal baseline for multi-hop networks, back-pressure scheduling, achieves its guarantee by propagating queue and channel state across the entire tree every slot [19]. The cost of this information is rarely accounted for, yet a centralized scheme that gathers global state in each slot scales poorly as the network grows, and the unobservable backhaul rate that a local scheme would need is exactly the quantity such global state was supplying.

Fifth, the uplink direction is frequently omitted, even though backhaul and access carry traffic in both directions and the self-interference coupling constrains uplink and downlink jointly at every node.

These gaps motivate a scheduler for a sectorized full-duplex multi-hop IAB network in which residual self-interference is modeled, the cost of global information is taken into account, and both traffic directions are scheduled.

1.2 Contributions

This thesis develops a scheduling framework for FD-IAB networks comprising a global-information policy, a local policy, and a stability analysis of both. The main contributions are as follows.

Distributed dynamic programming scheduler. We develop a distributed throughput-optimal scheduler for FD-IAB that accounts for the self-interference each receive link suffers. We solve it with a forward-backward dynamic programming algorithm that returns the optimal schedule in very low complexity, almost linear in the number of links.

Local scheduler. We design FD-LocalMW-BP, a local scheduler that replaces the exchange of global state with estimated backhaul rate based on a global policy. We show that the proposed policy achieves most of the stability region.

Empirical evaluation. A numerical study reveals that full-duplex scheduling sustains a far larger load than half-duplex. The proposed local policy attains a delay close to that of schedulers that use global information, while relying only on local information.

1.3 Organization

The remainder of this thesis is organized as follows. Chapter 2 reviews the technical background, including full-duplex operation, the per-link rate model under residual self-interference, and the queueing-theoretic preliminaries that underlie the analysis. Chapter 3 develops the FD-IAB scheduling framework. It formulates the per-slot maximum-weight problem, presents the dynamic programming solution, introduces the local scheduler FD-LocalMW-BP with its estimation phase, and establishes stability of both policies through fluid-model arguments. Chapter 4 reports numerical experiments comparing the schedulers across arrival rate, tree depth, and self-interference. Chapter 5 concludes and discusses directions for future work.

Chapter 2

Background

This chapter reviews the technical background for the rest of the dissertation. Section 2.1 introduces full-duplex operation, the self-interference it creates, and the per-link rate model that captures residual self-interference after cancellation. Section 2.2 collects the queueing-theoretic preliminaries used in the stability analysis, namely the notions of stability and stability region, the Lyapunov drift method, the Foster-Lyapunov criterion, and the MaxWeight policy.

2.1 Full-Duplex Operation and Self-Interference

A wireless transceiver operates in full-duplex (FD) mode when it transmits and receives on the same frequency band at the same time, in contrast to half-duplex (HD) mode where transmission and reception are separated in time or frequency. An idealized FD link doubles the spectral efficiency of an HD link by carrying both directions over the same radio resources.

The obstacle to realizing this gain is self-interference (SI), the leakage of the transmit signal of a node into its own colocated receive chain. Because the transmit and receive antennas are separated by only centimeters, the SI power can exceed the desired signal from a remote node by many orders of magnitude. Measurements at 28 GHz report interference-to-noise ratios from 0 to 40 dB, and worst-case values in practical transceivers can reach 90 dB when no mitigation is applied [14, 15]. Self-interference cancellation (SIC) suppresses this leakage

through a combination of passive isolation, analog cancellation, and digital cancellation, and at mmWave frequencies the directionality of large arrays adds a spatial degree of freedom, since transmit and receive beams can be shaped to be nearly orthogonal through the SI channel [15]. These techniques reduce SI by tens of dB but never remove it entirely. Hardware nonlinearities, finite receiver dynamic range, and channel-estimation error leave a residual SI component, so the realizable benefit of FD depends on how this residual is managed rather than on its idealized removal.

2.1.1 Rate Model under Residual Self-Interference

The IAB node studied in this dissertation hosts multiple antenna panels, several of which may transmit while another panel on the same node receives. Each transmit panel couples into the receive panel through its own SI channel, and because the transmit streams across panels are independent and zero-mean, the residual SI power at the receiver is the sum of the per-panel contributions. Normalizing by the noise power and absorbing the geometry and beam-design constants of each panel into a single per-panel quantity INR_ℓ , the aggregate interference-to-noise ratio at the receiver of link ℓ scales linearly with the number of co-active transmit panels,

$$\text{INR}_\ell^{\text{total}} = m_\ell \cdot \text{INR}_\ell, \quad (2.1)$$

where m_ℓ is the number of transmit-mode panels active alongside the receive panel on the same node. This linear scaling is the foundation of the per-link rate model used throughout the dissertation. Each additional co-active transmit panel adds one unit of INR_ℓ to the SINR denominator, so the instantaneous rate $R_\ell(\mathbf{s}, t)$ of a receive link decreases in m_ℓ , and the scheduling problem is the task of choosing how many panels to activate in transmit given this penalty.

2.2 Preliminaries

This section collects the queueing-theoretic results used in the stability analysis of Chapter 3. We follow the standard treatment of [7, 11, 19]. Let $\mathbf{q}(t) = (q_1(t), \dots, q_F(t))$ be the queue state of a network with F flows, evolving as $q_f(t+1) = (q_f(t) - \mu_f^\pi(t) + \lambda_f(t))^+$ under a scheduling policy π , where $\lambda_f(t)$ is the arrival process and $\mu_f^\pi(t)$ the offered service.

Definition 1 (Stability). A flow with queue length $q(t)$ is strongly stable under a policy π if

$$\limsup_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}^\pi[q(t)] < \infty. \quad (2.2)$$

A network is stable under π if every flow in it is strongly stable, equivalently if the Markov chain on the joint queue state $\mathbf{q}(t)$ is positive recurrent with a stationary distribution of finite first moment.

Strong stability implies a bounded long-run average queue and, by Little's law, a bounded long-run average delay.

Definition 2 (Stability region). The stability region $\mathcal{C}$ of a network is the set of arrival-rate vectors for which some scheduling policy stabilizes the network,

$$\mathcal{C} := \{ \boldsymbol{\lambda} : \exists \pi \text{ stabilizing under } \boldsymbol{\lambda} \}. \quad (2.3)$$

A policy is throughput-optimal if it stabilizes every $\boldsymbol{\lambda}$ in the interior of $\mathcal{C}$.

The Lyapunov drift method reduces a stability question to a per-slot inequality on the queues. For the quadratic Lyapunov function

$$V(\mathbf{q}) = \frac{1}{2} \sum_{f=1}^F q_f^2, \quad (2.4)$$

which is non-negative and zero only at the empty state, the one-step drift under π is

$$\Delta V(\mathbf{q}) := \mathbb{E}^\pi[V(\mathbf{q}(t+1)) - V(\mathbf{q}(t)) \mid \mathbf{q}(t) = \mathbf{q}]. \quad (2.5)$$

Expanding the quadratic and substituting the queue update yields the standard bound

$$\Delta V(\mathbf{q}) \leq B - \sum_f q_f (\mu_f^\pi - \lambda_f), \quad (2.6)$$

where μ_f^π is the expected service rate of flow f under π and B is a constant bounded by the second moments of the arrival and service processes.

Theorem 1 (Foster-Lyapunov criterion). *Let $\{\mathbf{q}(t)\}$ be a Markov chain and V a non-negative function. If there exist a finite set $\mathcal{S}_0$ and constants $\epsilon, B' > 0$ such that*

$$\Delta V(\mathbf{q}) \leq -\epsilon \quad \text{for all } \mathbf{q} \notin \mathcal{S}_0, \quad (2.7)$$

and $\Delta V(\mathbf{q}) \leq B'$ for $\mathbf{q} \in \mathcal{S}_0$, then the chain is positive recurrent [10, Theorem 14.1.5].

Combining Theorem 1 with (2.6) reduces stability to a per-flow rate condition. If the offered service satisfies $\mu_f^\pi > \lambda_f$ for every flow, the negative drift term eventually dominates the constant B and the network is stable.

Definition 3 (MaxWeight policy). In a single-hop network, the MaxWeight policy chooses, in each slot, the feasible service-rate vector that minimizes the right-hand side of (2.6) given the current queues,

$$\pi^{\text{MW}}(\mathbf{q}) = \arg \max_{\boldsymbol{\mu} \in \text{Conv}(\mathcal{S})} \sum_f q_f \cdot \mu_f, \quad (2.8)$$

where $\mathcal{S}$ is the set of feasible service-rate vectors and $\text{Conv}(\cdot)$ its convex hull.

Theorem 2 (Throughput optimality of MaxWeight). *The MaxWeight policy of Definition 3 is throughput-optimal. It stabilizes every arrival-rate vector in the interior of the stability region $\mathcal{C}$ [19].*

Table 2.1: Notation: topology, nodes, panels, time.

Symbol	Meaning
t	Discrete time slot index
T	Simulation horizon (slots)
T_s	Slot duration (seconds)
BW	Bandwidth (Hz)
PKT	Packet size (bits)
$\mathcal{G}$	Directed tree representing the IAB network
$\mathcal{B}$	Set of IAB nodes (gNBs)
$\mathcal{U}$	Set of User Equipments (UEs)
r	Donor gNB (root of $\mathcal{G}$)
N	Total number of IAB nodes ($N = \mathcal{B} $)
n	IAB node index
$\text{pr}(n)$	Parent IAB of node n
$\mathcal{C}(n)$	Set of children of node n
$\mathcal{T}_n$	Subtree rooted at n
P_n	Number of panels at IAB n
$\mathcal{P}_n$	Panel set of IAB n , $\{P_n^1, \dots, P_n^{P_n}\}$
P_n^1	Panel of n serving parent-backhaul link b_n
$\mathcal{P}$	Set of all panels in the network, $\bigcup_{n \in \mathcal{B}} \mathcal{P}_n$
$I(p)$	IAB node hosting panel p
k	Panel index

The FD-LocalMW policy developed in Chapter 3 builds on this framework. It replaces the unobservable backhaul-downlink rate in the MaxWeight rule with an estimated weight and applies the same drift argument in the resulting partially observed setting.

2.3 Notation and Conventions

Throughout the dissertation, we use the conventions summarized in Tables 2.1–2.4:

Vectors are bold-face: $\mathbf{q}, \boldsymbol{\lambda}$. The positive-part operator $(x)^+ := \max(0, x)$ is used throughout. Expectations are with respect to the joint distribution of arrivals, channel fading, and any randomization in the policy. Where disambiguation is needed, $\mathbb{E}^\pi$ denotes expectation under policy π .

Table 2.2: Notation: links, scheduling, capacity.

Symbol	Meaning
$\mathcal{L}$	Link set of the network
ℓ	Link index
$A(\ell), D(\ell)$	Source / destination of link ℓ
$\mathcal{L}_{\text{UL}}, \mathcal{L}_{\text{DL}}$	Uplink / downlink links
$\mathcal{L}_{\text{AC}}, \mathcal{L}_{\text{BH}}$	Access / backhaul links
b_n	Parent-backhaul link of IAB n (from $\text{pr}(n)$ to n)
$\mathcal{L}_n$	Links incident to a panel of IAB n
$\mathcal{L}_n^{\text{TX}}, \mathcal{L}_n^{\text{RX}}$	TX-role / RX-role links at IAB n
$\mathcal{L}_n^{\text{AC}}$	Access links incident to IAB n
$\mathcal{L}_{n,k}$	Links hostable on panel k of IAB n
$s_\ell(t)$	Schedule indicator of link ℓ in slot t
$\mathbf{s}(t)$	Activation vector $(s_\ell(t))_{\ell \in \mathcal{L}}$
$\mathcal{S}$	Set of feasible activation vectors
$m_\ell(\mathbf{s}, t)$	SI count: TX panels active at $I(D(\ell))$ in slot t
SNR_ℓ	Signal-to-noise ratio on link ℓ
INR_ℓ	Per-panel interference-to-noise ratio at receiver of ℓ
$R_\ell(\mathbf{s}, t)$	Instantaneous capacity of ℓ (packets/slot)
$\boldsymbol{\mu}, \mu_\ell$	Channel state vector, channel rate on link ℓ
$\pi_{\boldsymbol{\mu}}$	Stationary distribution of $\boldsymbol{\mu}$

Table 2.3: Notation: flows, queues, weights, DP messages.

Symbol	Meaning
$\mathcal{F}$	Flow set
f	Flow index
$\text{src}(f), \text{dst}(f)$	Source / destination node of flow f
$\mathcal{F}_\ell$	Flows traversing link ℓ
$\Lambda_f^{\text{ext}}(t)$	Exogenous arrivals to flow f in slot t
λ_f^{ext}	Mean exogenous arrival rate of flow f
$\boldsymbol{\nu}, \nu_\ell$	Per-link arrival rate vector and its components
$q_n^f(t)$	Queue length of flow f at node n , slot t
$d_n^f(t)$	Flow- f packets departing node n in slot t
$Q_\ell(t)$	Aggregate per-link queue, $\sum_{f \in \mathcal{F}_\ell} q_{A(\ell)}^f(t)$
w_ℓ	Weight of access-DL link ℓ
$\bar{w}_\ell(m)$	Weight of access-UL link ℓ at SI level m
$w_{b_n}^{\text{DL}}(m)$	Weight of BH-DL on b_n at SI level m
$w_{b_n}^{\text{UL}}(m_{\text{pr}(n)})$	Weight of BH-UL on b_n at SI level $m_{\text{pr}(n)}$ of the parent
$\nu_{n,\ell}^{\text{T}}$	DP message: best subtree weight when n transmits on ℓ
$\nu_{n,\ell}^{\text{R}}(m)$	DP message: best subtree weight when n receives on ℓ at SI level m
$\nu_{n,\ell}^{\text{IDLE}}$	DP message: best subtree weight when n turns off ℓ
$\nu_{n,k}^{\text{T}}$	DP message at panel level, panel k in TX mode
$\nu_{n,k}^{\text{R}}(m)$	DP message at panel level, panel k in RX mode at SI level m
$\nu_{n,k}^{\text{IDLE}}$	DP message at panel level, panel k idle
ν_n^{T}	DP message at node level, $\text{pr}(n)$ transmits to n on b_n
ν_n^{R}	DP message at node level, $\text{pr}(n)$ receives from n on b_n
ν_n^{IDLE}	DP message at node level, b_n idle

Table 2.4: Notation: stability, modes, local policy, fluid model.

Symbol	Meaning
Λ	Global stability region
Λ^{loc}	Local stability region (achievable by local policies)
θ_n	Mode of node n : T - m (BH-DL RX with m co-TX) or ϕ (BH idle)
$\mathcal{M}_n$	Mode set of node n
$\mathcal{S}_n(\theta_n)$	Per-mode feasible activation set at n
$\mathcal{S}_n^{T,m}, \mathcal{S}_n^{\text{IDLE}}$	Feasible sets in Mode T - m and Mode ϕ
$R_{b_n}(\theta_n, \boldsymbol{\mu}_n)$	BH-DL rate of b_n in mode θ_n , channel $\boldsymbol{\mu}_n$
$\hat{\mu}_{b_c}$	Estimated BH-DL weight (conditional mean rate under BP)
ρ_n	Total Mode T fraction at node n
$\alpha_n(m, \boldsymbol{\mu}_n)$	Joint mode-channel fraction (LP variable)
$y_{n,\ell}^{T,m}(\boldsymbol{\mu}_n)$	Within-mode activation fraction in Mode T - m (LP variable)
$y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n)$	Within-mode activation fraction in Mode ϕ (LP variable)
τ	Continuous time in the fluid model
$\bar{q}_\ell(\tau)$	Fluid-scaled queue length
ϵ_{train}	Training estimation error of $\hat{\mu}$
g_n	Per-node information gap
$\epsilon^*(\boldsymbol{\nu})$	Critical shrinkage of FD-LocalMW at reference load $\boldsymbol{\nu}$
$V(\mathbf{q})$	Lyapunov function
ΔV	Lyapunov drift
$(\cdot)^+$	$\max(0, \cdot)$

Chapter 3

Full-duplex Integrated Access Backhaul Scheduling

3.1 System Model

3.1.1 Tree, Nodes, Panels

We consider a full-duplex IAB network represented by a directed tree $\mathcal{G} = (\mathcal{B} \cup \mathcal{U}, \mathcal{L}, r)$, where $\mathcal{B} = \{1, \dots, |\mathcal{B}|\}$ is the set of IAB nodes (gNBs), $\mathcal{U} = \{1, \dots, |\mathcal{U}|\}$ is the set of User Equipments (UEs), and $r \in \mathcal{B}$ is the donor gNB at the root. For every non-root IAB $b \in \mathcal{B} \setminus \{r\}$, $\text{pr}(b)$ denotes its parent and $\text{ch}(b) \subset \mathcal{B} \cup \mathcal{U}$ its children. Each IAB $i \in \mathcal{B}$ carries a fixed set of analog-beamforming panels

$$\mathcal{P}_i := \{p_{i,1}, p_{i,2}, \dots, p_{i,|\mathcal{P}_i|}\}, \quad (3.1)$$

and we write $\mathcal{P} := \bigcup_{i \in \mathcal{B}} \mathcal{P}_i$ for the set of all panels in the network. The map $I : \mathcal{P} \rightarrow \mathcal{B}$ returns the IAB hosting a given panel, so $p \in \mathcal{P}_{I(p)}$ by construction. We assume that all panels within a single IAB are physically identical and have the same SI level. Due to the high directivity of the mmWave beams used at access frequencies, we further assume that the

Class	Endpoints ($A(\ell) \rightarrow D(\ell)$)	SI at receiver?
$\mathcal{L}_{\text{AC}} \cap \mathcal{L}_{\text{DL}}$	IAB panel $\rightarrow$ UE	no
$\mathcal{L}_{\text{AC}} \cap \mathcal{L}_{\text{UL}}$	UE $\rightarrow$ IAB panel	yes
$\mathcal{L}_{\text{BH}} \cap \mathcal{L}_{\text{DL}}$	parent IAB panel $\rightarrow$ child IAB panel	yes
$\mathcal{L}_{\text{BH}} \cap \mathcal{L}_{\text{UL}}$	child IAB panel $\rightarrow$ parent IAB panel	yes

Table 3.1: Four physical link classes induced by the joint direction/access partition of $\mathcal{L}$. Only the AC-DL class has a UE receiver and is therefore SI-free; the other three classes terminate at an IAB panel and are exposed to in-IAB SI.

coverage area of any panel is disjoint from that of every other panel in the network, so that intended-link interference between distinct panels can be neglected.

3.1.2 Links and their Orientation

A link is an ordered pair $\ell : A(\ell) \rightarrow D(\ell)$ with source node $A(\ell)$ and destination node $D(\ell)$, where $A(\ell), D(\ell) \in \mathcal{U} \cup \mathcal{P}$. The source emits a beamformed signal and the destination receives it. The full link set $\mathcal{L}$ admits two independent partitions,

$$\mathcal{L} = \mathcal{L}_{\text{UL}} \cup \mathcal{L}_{\text{DL}} \quad (\text{direction}), \quad (3.2)$$

$$= \mathcal{L}_{\text{AC}} \cup \mathcal{L}_{\text{BH}} \quad (\text{access and backhaul}), \quad (3.3)$$

the first separating uplink from downlink user-plane traffic and the second separating access links (one UE endpoint) from backhaul links (both endpoints are IAB panels). In the tree, downlink links point away from the root and uplink links point toward the root. The four cells of the joint partition correspond to the four physical link classes summarized in Table 3.1.

3.1.3 Link Decomposition

For an IAB $i \in \mathcal{B}$, let $\mathcal{L}_i \subseteq \mathcal{L}$ denote the subset of links incident on a panel of i . We further split $\mathcal{L}_i$ by the role i plays in the link:

$$\begin{aligned} \mathcal{L}_i &= \mathcal{L}_i^{\text{TX}} \cup \mathcal{L}_i^{\text{RX}}, & \mathcal{L}_i^{\text{TX}} &:= \{\ell \in \mathcal{L}_i : A(\ell) \in \mathcal{P}_i\}, \\ & & \mathcal{L}_i^{\text{RX}} &:= \{\ell \in \mathcal{L}_i : D(\ell) \in \mathcal{P}_i\}. \end{aligned} \quad (3.4)$$

A link ℓ belongs to $\mathcal{L}_i^{\text{TX}}$ if it is *transmitted by* a panel of i , and to $\mathcal{L}_i^{\text{RX}}$ if it is *received by* a panel of i . This decomposition determines the SI geometry: a panel in RX mode at IAB i experiences self-interference only from the concurrently active TX-mode panels of the same IAB i .

3.1.4 Scheduling Variable and the SI count

We adopt a slotted-time model with slots indexed by $t \in \mathbb{Z}_+$. At each slot, the scheduler selects a binary activation

$$s_\ell(t) \in \{0, 1\}, \quad \ell \in \mathcal{L}, \quad (3.5)$$

where $s_\ell(t) = 1$ indicates that link ℓ is scheduled in slot t . We collect $\mathbf{s}(t) := (s_\ell(t))_{\ell \in \mathcal{L}}$.

The *SI count* m_ℓ is the number of TX-mode panels at the receiver IAB of link ℓ that are simultaneously active:

$$m_\ell(\mathbf{s}, t) := \sum_{\ell' \in \mathcal{L}_{I(D(\ell))}^{\text{TX}}} s_{\ell'}(t), \quad \ell \in \mathcal{L} \setminus (\mathcal{L}_{\text{AC}} \cap \mathcal{L}_{\text{DL}}), \quad (3.6)$$

On access-downlink links $\ell \in \mathcal{L}_{\text{AC}} \cap \mathcal{L}_{\text{DL}}$, the receiver $D(\ell)$ is a UE (no co-located TX panels), so $m_\ell \equiv 0$. When node n operates in Mode T with TX count m , the BH-DL link b_n has $m_{b_n} = m$.

3.1.5 Link Capacity

Given the SI count $m_\ell(\mathbf{s}, t)$, the instantaneous capacity of link ℓ in slot t is

$$R_\ell(\mathbf{s}, t) = \frac{\text{BW} \cdot T_s}{\text{PKT}} \log_2 \left(1 + \frac{\text{SNR}_\ell(t)}{1 + m_\ell(\mathbf{s}, t) \cdot \text{INR}_\ell} \right), \quad (3.7)$$

in packets per slot, where BW is the bandwidth, T_s the slot duration, PKT the packet size in bits, SNR_ℓ the signal-to-noise ratio on link ℓ , and INR_ℓ the per-panel interference-to-noise ratio at the receiver of ℓ .

3.1.6 Flows and Queueing Dynamics

Flows Traffic in the network is organized into *flows*. A flow $f \in \mathcal{F}$ is the ordered path of nodes connecting a source $\text{src}(f)$ to a destination $\text{dst}(f)$ in the tree. The flow set partitions by direction,

$$\mathcal{F} = \mathcal{F}_{\text{UL}} \cup \mathcal{F}_{\text{DL}}, \quad (3.8)$$

where uplink flows $f \in \mathcal{F}_{\text{UL}}$ originate at a UE ($\text{src}(f) \in \mathcal{U}$) and terminate at the donor ($\text{dst}(f) = r$), and downlink flows $f \in \mathcal{F}_{\text{DL}}$ reverse the roles ($\text{src}(f) = r$, $\text{dst}(f) \in \mathcal{U}$). Because $\mathcal{G}$ is a tree, each flow's path is uniquely determined by its endpoints. We write $\ell \in f$ when link ℓ lies on flow f 's path, and define $\mathcal{F}_\ell := \{f \in \mathcal{F} : \ell \in f\}$ as the set of flows traversing link ℓ .

Queues Each node $n \in \mathcal{B} \cup \mathcal{U}$ participating in flow f maintains a separate buffer with instantaneous queue length $q_n^f(t)$, with $q_n^f(t) \equiv 0$ for $n \notin f$. Packets arrive exogenously only at the source of each flow: the external arrival to flow f at $\text{src}(f)$ is

$$\Lambda_f^{\text{ext}}(t) \stackrel{\text{i.i.d.}}{\sim} \text{Poisson}(\lambda_f^{\text{ext}}), \quad \lambda_f^{\text{ext}} := \mathbb{E}[\Lambda_f^{\text{ext}}(t)]. \quad (3.9)$$

We denote by $d_n^f(t) \in \mathbb{Z}_+$ the number of flow- f packets that depart node n in slot t . Since each node has a unique outgoing link on flow f 's path, $d_n^f(t)$ is delivered on exactly that link and is jointly determined by the activation $s_\ell(t)$ of the outgoing link, its link capacity $R_\ell(\mathbf{s}, t)$ from (3.7). The aggregate served on link ℓ obeys the capacity constraint

$$\sum_{f \in \mathcal{F}_\ell} d_{A(\ell)}^f(t) \leq R_\ell(\mathbf{s}, t) \cdot s_\ell(t), \quad (3.10)$$

which couples queue evolution to the activation vector $\mathbf{s}(t)$ defined in Section 3.1.1.

Queue evolution For the root r and a non-root IAB $b \in \mathcal{B} \setminus \{r\}$, the per-flow queue update is

$$q_r^f(t+1) = (q_r^f(t) - d_r^f(t))^+ + \begin{cases} \sum_{c \in \text{ch}(r) \cap f} d_c^f(t), & f \in \mathcal{F}_{\text{UL}}, \\ \Lambda_f^{\text{ext}}(t), & f \in \mathcal{F}_{\text{DL}}, \end{cases} \quad (3.11)$$

$$q_b^f(t+1) = (q_b^f(t) - d_b^f(t))^+ + \begin{cases} \sum_{c \in \text{ch}(b) \cap f} d_c^f(t), & f \in \mathcal{F}_{\text{UL}}, \\ d_{\text{pr}(b)}^f(t), & f \in \mathcal{F}_{\text{DL}}. \end{cases} \quad (3.12)$$

The negative term $d_n^f(t)$ in each equation is the service drained from node n on its outgoing link in slot t , while the positive term collects the inflow: an external Poisson arrival when n is the flow's source, or the service that completed in slot t on the incoming link from n 's parent (downlink case) or its child on f 's path (uplink case). The sum-over-children in the uplink case contains at most one nonzero term, since f traverses a unique child of n . A UE $u \in \mathcal{U}$ at the source of an uplink flow obeys the recursion of (3.12) with the sum-of-children term replaced by $\Lambda_f^{\text{ext}}(t)$. The projection $(\cdot)^+ := \max(\cdot, 0)$ guards against the scheduler over-allocating service on an empty queue.

Aggregate Queue The scheduling and stability analysis of Chapter 3 operates on the *aggregate per-link queue*, defined for each link ℓ as the total backlog at its source node summed across the flows that traverse ℓ :

$$Q_\ell(t) := \sum_{f \in \mathcal{F}_\ell} q_{A(\ell)}^f(t). \quad (3.13)$$

$Q_\ell(t)$ is the workload that link ℓ is competing to serve in slot t . And define the arrival rate per flow as:

$$\lambda_\ell := \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[\Lambda_\ell(t)] = \sum_{f \in \mathcal{F}_\ell} \lambda_f^{\text{ext}}, \quad (3.14)$$

Maximum Weighted Schedule in IAB Network A wide class of scheduling policies for queueing networks reduces, in each slot, to a maximum-weight optimization over the activation vector. Given a per-link weight $w_\ell(t)$ on every link $\ell \in \mathcal{L}$, the schedule is chosen by

$$\mathbf{s}(t) \in \arg \max_{\mathbf{s} \in \mathcal{S}} \sum_{\ell \in \mathcal{L}} w_\ell(t) s_\ell(t), \quad (3.15)$$

where $\mathcal{S} \subseteq \{0, 1\}^{|\mathcal{L}|}$ is the set of activation vectors satisfying the per-panel feasibility constraint of Section 3.1.1. [19] introduced the MaxWeight framework, in which the schedule maximizes the weighted sum (3.15) each slot. Their original formulation uses the per-flow differential backlog $w_\ell(t) := R_\ell(\mathbf{s}, t) \max_{f \in \mathcal{F}_\ell} (q_{\text{pr}(n)}^f(t) - q_n^f(t))$ and is throughput-optimal for multi-hop networks, the policy commonly called *back-pressure*. A simpler variant replaces the differential with the aggregate per-link queue, taking $w_\ell(t) := Q_\ell(t) R_\ell(\mathbf{s}, t)$. We refer to the former as the *back-pressure policy* (Section 3.1.8) and the latter as the *MaxWeight policy* (Section 3.1.9).

A central feature of the FD-IAB instance of (3.15) is that the weight on link ℓ is not given a priori. Through the capacity $R_\ell(\mathbf{s}, t)$, it depends on the schedule only through the transmit count $m_\ell(t)$ of (3.6), the number of transmitting panels at the IAB receiving ℓ ,

$$w_\ell(t) = w_\ell(m_\ell(t), t). \quad (3.16)$$

Switching one more panel of that IAB from idle to transmit raises m_ℓ , lowers R_ℓ , and so changes w_ℓ , hence the objective in (3.15) is not separable over the links incident on a single IAB.

Solving (3.15) is the computational core of every scheduling policy in this class. In this chapter, we develop a forward-backward dynamic-programming algorithm that obtains the maximum-weighted schedule in $\mathcal{O}(|\mathcal{L}|)$ time despite the within-IAB weight coupling described above.

A dynamic-programming approach of similar flavor is known for finding a maximum-weighted independent set on a tree [4], but two structural features of the FD-IAB problem set our instance apart from that classical recursion. The first is the multi-panel structure: at each IAB i , up to $|\mathcal{P}_i|$ panels may be active concurrently, so the DP sub-problem at the node is not a binary include/exclude choice over a single vertex but the selection of a subset of links of cardinality at most $|\mathcal{P}_i|$ from a candidate set $\mathcal{L}_i$. The second is the schedule-dependent weight (3.16) introduced by full-duplex operation: even after fixing which subset of links to activate, the contribution of each link to the objective must be recomputed under the SI count induced by the subset itself, since no link's w_ℓ is fixed until $\mathbf{s}_{-\ell}$ is fixed. The two effects together turn the per-IAB DP sub-problem into a combinatorial optimization with schedule-dependent costs, for which we provide an optimal greedy algorithm with running time almost linear in $|\mathcal{P}_i|$ (Section 3.1.6).

Given the link weights, the weight $w_\ell(m_\ell, t)$ of each link depends on the schedule only through the SI count m_ℓ of (3.6), the number of transmitting panels at the IAB receiving ℓ .

For every non-leaf node $n \in \mathcal{B} \cup \mathcal{U}$, let $\mathcal{G}_n$ be the sub-tree of $\mathcal{G}$ induced by n and its

descendants and let $\mathcal{L}_n$ be its links. The maximum-weight program restricted to $\mathcal{G}_n$ is

$$\max_{\mathbf{s}} \sum_{\ell \in \mathcal{L}_n} w_\ell(m_\ell, t) s_\ell \quad (3.17)$$

$$\text{s.t. } \sum_{\ell \in \mathcal{L}_p} s_\ell \leq 1, \quad \forall p \in \mathcal{P}_i, \forall i \in \mathcal{B} \cap \mathcal{G}_n, \quad (3.18)$$

$$s_\ell \in \{0, 1\}, \quad \forall \ell \in \mathcal{L}_n, \quad (3.19)$$

where $\mathcal{L}_p := \{\ell \in \mathcal{L} : A(\ell) = p \text{ or } D(\ell) = p\}$ is the set of links touching panel p .

Flows, Routes, and Arrival Rates Let $\mathcal{F}$ denote the set of flows. Each flow $f \in \mathcal{F}$ originates at the root (donor) and is destined for a specific UE served by access link $\ell_f \in \mathcal{A}$ at some node n_f . The flow's route $\text{path}(f)$ is the unique sequence of links from root to ℓ_f through the tree:

$$\text{path}(f) = (b_{n_1}, b_{n_2}, \dots, b_{n_f}, \ell_f), \quad (3.20)$$

where $b_{n_1}, \dots, b_{n_f}$ are the BH-DL links traversed from root to node n_f , and ℓ_f is the final access link. Since the topology is a tree, each flow's route is unique.

External packets for flow f arrive at the root as an i.i.d. process with mean ν^f packets per slot and finite second moment:

$$\mathbb{E}[A^f(t)] = \nu^f, \quad \mathbb{E}[(A^f(t))^2] < \infty, \quad \text{i.i.d. across } t, \quad \text{independent across } f. \quad (3.21)$$

The *flow arrival rate vector* $\boldsymbol{\nu} := (\nu^f)_{f \in \mathcal{F}}$ parametrizes the traffic demand. The per-link arrival rates are induced by $\boldsymbol{\nu}$:

$$\nu_\ell := \sum_{f \in \mathcal{F}: \ell_f = \ell} \nu^f, \quad \text{for each access link } \ell \in \mathcal{A}, \quad (3.22)$$

$$\nu_{b_n} := \sum_{f \in \mathcal{F}: b_n \in \text{path}(f)} \nu^f, \quad \text{for each BH-DL link } b_n. \quad (3.23)$$

By the tree structure, $\nu_{b_n} = \sum_{\ell \in \text{subtree}(n)} \nu_\ell$: the BH-DL load at node n equals the total access demand in its subtree. Each BH-DL queue q_{b_n} at the parent carries packets from multiple flows: $q_{b_n} = \sum_{\ell \in \text{subtree}(n)} q_{b_n}^{(\ell)}$, where $q_{b_n}^{(\ell)}$ is the share destined for access link ℓ .

3.1.7 Stability Region

We now characterize the set of arrival rate vectors for which the network can be stabilized. Let $\mathcal{M}$ denote the set of all possible channel state vectors $\boldsymbol{\mu} = (\mu_\ell)_{\ell \in \mathcal{L}}$, and let $\pi_\mu := \Pr[\boldsymbol{\mu}(t) = \boldsymbol{\mu}]$ be the stationary probability of state $\boldsymbol{\mu}$.

For a given channel state $\boldsymbol{\mu} \in \mathcal{M}$ and a feasible schedule $\mathbf{s} \in \mathcal{S}$, the rate delivered on link ℓ is $R_\ell(\mathbf{s}) \cdot s_\ell$, where $R_\ell(\mathbf{s})$ is the capacity from (3.7) evaluated at the SI count $m_\ell(\mathbf{s})$ induced by $\mathbf{s}$. The set of rate vectors achievable under channel state $\boldsymbol{\mu}$ is

$$\mathcal{C}_\mu := \left\{ (R_\ell(\mathbf{s}) \cdot s_\ell)_{\ell \in \mathcal{L}} : \mathbf{s} \in \mathcal{S} \right\}. \quad (3.24)$$

The stability region is defined as

$$\Lambda := \left\{ \boldsymbol{\nu} = [\nu^f]_{f \in \mathcal{F}} : [\nu_\ell]_{\ell \in \mathcal{L}} \in \sum_{\boldsymbol{\mu} \in \mathcal{M}} \pi_\mu \text{Conv}(\mathcal{C}_\mu) \right\}, \quad (3.25)$$

where $\nu_\ell := \sum_{f \in \mathcal{F}_\ell} \nu^f$ is the aggregate arrival rate on link ℓ , and $\text{Conv}(\cdot)$ denotes the convex hull. A rate vector $\boldsymbol{\nu}$ is in the interior of Λ if there exists $\delta > 0$ such that $\boldsymbol{\nu} + \delta \mathbf{1} \in \Lambda$.

Definition 4. The system is *stable* under a scheduling policy if and only if $\limsup_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[\|\mathbf{Q}(t)\|] < \infty$, where $\mathbf{Q}(t) := (Q_\ell(t))_{\ell \in \mathcal{L}}$. The system is *stabilizable* if there exists a stationary scheduling policy under which it is stable.

Lemma 1. *If $\boldsymbol{\nu} \notin \Lambda$, the network is not stabilizable under any policy. If $\boldsymbol{\nu}$ is in the interior of Λ , there exists a stationary scheduling policy that stabilizes the network.*

Proof. The proof follows standard arguments from [11, 19]. □

3.1.8 Global Back-Pressure Policy

We now define a scheduling policy that stabilizes the network for any arrival rate vector in the interior of Λ . The construction follows the framework of [19].

For an access link ℓ incident to node n (i.e., $\ell \in \mathcal{L}_{AC} \cap \mathcal{L}_n$), define the back-pressure weight

$$w_\ell^{\text{bp}}(t) := R_\ell(\mathbf{s}, t) \max_{f \in \mathcal{F}_\ell} q_n^f(t). \quad (3.26)$$

For a backhaul link b_n connecting $\text{pr}(n)$ to n , define

$$w_{b_n}^{\text{bp}}(t) := R_{b_n}(\mathbf{s}, t) \max_{f \in \mathcal{F}_{b_n}} \left(q_{\text{pr}(n)}^f(t) - q_n^f(t) \right). \quad (3.27)$$

In each slot, the back-pressure policy selects

$$\mathbf{s}^{\text{bp}}(t) \in \arg \max_{\mathbf{s} \in \mathcal{S}} \sum_{\ell \in \mathcal{L}} w_\ell^{\text{bp}}(t) s_\ell(t). \quad (3.28)$$

When a backhaul link b_n is scheduled, only the packets of the flow $f^\star := \arg \max_{f \in \mathcal{F}_{b_n}} (q_{\text{pr}(n)}^f(t) - q_n^f(t))$ are transmitted.

Since $R_\ell(\mathbf{s}, t)$ depends on the SI count $m_\ell(\mathbf{s}, t)$, the capacity of each link is itself a function of the schedule being optimized. The per-slot problem (3.28) is therefore an instance of the MWS problem (3.17) with schedule-dependent weights.

Lemma 2. *The back-pressure policy (3.28) stabilizes the network for any arrival rate vector $\boldsymbol{\nu}$ in the interior of Λ .*

Proof. Throughput-optimality follows from the standard per-flow Lyapunov function $L(t) := \frac{1}{2} \sum_{n \in \mathcal{B}} \sum_{f \in \mathcal{F}} (q_n^f(t))^2$ and use standard proof of [11, 19]. $\square$

3.1.9 Global Max-Weight Policy

A simpler variant replaces the per-flow differential backlog with the aggregate per-link queue $Q_\ell(t)$. The max-weight policy selects

$$\mathbf{s}^{\text{mw}}(t) \in \arg \max_{\mathbf{s} \in \mathcal{S}} \sum_{\ell \in \mathcal{L}} Q_\ell(t) R_\ell(\mathbf{s}, t) s_\ell(t). \quad (3.29)$$

Unlike back-pressure, the throughput optimality, which guarantees stability for all arrival vectors in the stability region, is not guaranteed. However, since the policy can transmit multiple packets at once over a single link, it achieves lower delay than back-pressure.

Both (3.28) and (3.29) are instances of the MWS problem (3.17): the weight $w_\ell(t)$ absorbs the queue information, and the per-slot optimization is solved by the DP algorithm developed in the next section. The only difference between the two policies is how the weights are computed from the queue state; the combinatorial structure of the per-slot optimization is identical.

3.2 Dynamic Programming approach

Solving a maximum weighted schedule, (3.17) is crucial for finding scheduling policies. In this section, we propose a forward-backward algorithm based on dynamic programming (DP) to find a maximum weighted schedule in almost linear time. The most naive approach for this problem is exhaustive search, which looks at every single feasible schedule. However, pragmatically, it is almost intractable to deploy the exhaustive method because it has exponential complexity, $O(2^L)$. Instead, the DP approach can make the complexity polynomial through the elimination of redundancy by passing messages with nodes. In existing half-duplex IAB schedulers, SI is avoided by prohibiting simultaneous transmission and reception, which makes the link weights independent of the schedule. However, in the FD system, the existing methodology can not be directly adopted since the SI makes the weights dependent on other

schedules. Furthermore, the UL traffic was overlooked by the work. Hence, a scheduling policy considering FD operations and UL-DL bidirectional flows is developed in this section. Through the optimal greedy algorithm with messages of subtree, the algorithm can achieve the almost linear complexity in the number of links.

To implement the algorithm, we will firstly derive the optimal sum of weights of the subtree from leaf to other cases.

3.2.1 Leaf Node

Let n be a node whose downlinks are all access links; $\mathcal{L}_n$ consists of access links and b_n and there are no child-BH links.

Definition 5. (Link level) For IAB n and link $\ell \in \mathcal{L}_n$, define $\nu_{n,\ell}^T, \nu_{n,\ell}^R(m), \nu_{n,\ell}^{\text{IDLE}}$ as the optimal weighted sum on the part of subtree of n and backhaul, $b_n, \mathcal{T}_n \cup \{b_n\}$. Then, on the leaf node n , except backhaul link b_n , they can be obtained as

$$\nu_{n,\ell}^T = w_\ell, \quad \nu_{n,\ell}^R(m) = \bar{w}_\ell(m), \quad \nu_{n,\ell}^{\text{IDLE}} = 0 \quad \text{for } \ell \in \mathcal{L}_n \setminus b_n \quad (3.30)$$

Definition 6. (Panel level) For IAB n and panel $k \in \mathcal{P}_n$, define $\nu_{n,k}^T, \nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}$ as the optimal weighted sum on the subtree connected with panel k when the panel k is TX, RX, and IDLE, respectively. Then, on the leaf node n , they can be computed as

$$\nu_{n,k}^T = \max_{\ell \in \mathcal{L}_{n,k}} \nu_{n,\ell}^T, \quad \nu_{n,k}^R(m) = \max_{\ell \in \mathcal{L}_{n,k}} \nu_{n,\ell}^R(m), \quad \nu_{n,k}^{\text{IDLE}} = 0 \quad \text{for } k \in \mathcal{P}_n. \quad (3.31)$$

This is because, if the panel is active, due to (3.18), the panel should pick the most highest weight from $\mathcal{L}_{n,k}$. If the panel is idle, $\nu_{n,k}^{\text{IDLE}} = 0$ since it should not pick any links.

Definition 7. (Node level) For IAB n , if the parent of n , $\text{pr}(n)$ transmits through b_n to n , define ν_n^T as the optimal weighted sum of $\mathcal{T}_n$. Analogously, when the $\text{pr}(n)$ receives through b_n from n and b_n is idle, define ν_n^R and ν_n^{IDLE} as the optimal weighted sum of $\mathcal{T}_n$, respectively.

Lemma 3. Given P_n^1 serves the backhaul link of n and $\text{pr}(n)$, and transmits packets to n , ν_n^T can be calculated as

$$\nu_n^T = \max_{0 \leq m \leq P_n - 1} \left[w_{b_n}^{\text{DL}}(m) + \sum_{k \neq P_n^1} \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m} \sum_{k \in S} \max_{\ell \in \mathcal{L}_{n,k}} w_\ell - \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ \right] \quad (3.32)$$

Proof. We will scan from $m = 0$ to $m = P_n - 1$ to consider the effect of SI. If the number of active TX panels is m , the weight of backhaul link connected to parent is $w_{b_n}^{\text{DL}}(m)$. Since the number of TX panel is m , other $P_n - m - 1$ panels except P_n^1 should be either RX or IDLE. Now, partition panels $\mathcal{P}_n \setminus P_n^1$ into the set of TX panels, $S, |S| = m$ and non-TX panels, $(\mathcal{P}_n \setminus P_n^1) \setminus S$. Hence, the optimal weight sum of panels except P_n^1 is

$$\sum_{k \in S} \nu_{n,k}^T + \sum_{k \notin S} \max(\nu_{n,k}^{\text{IDLE}}, \nu_{n,k}^{\text{R}}(m)) = \sum_{k \in S} \nu_{n,k}^T + \sum_{k \notin S} (\nu_{n,k}^{\text{R}}(m))^+ \quad (3.33)$$

$$= \sum_{k \in S} \nu_{n,k}^T + \sum_{k \in \mathcal{P}_n \setminus P_n^1} (\nu_{n,k}^{\text{R}}(m))^+ - \sum_{k \in S} (\nu_{n,k}^{\text{R}}(m))^+ \quad (3.34)$$

$$= \sum_{k \in \mathcal{P}_n \setminus P_n^1} (\nu_{n,k}^{\text{R}}(m))^+ + \sum_{k \in S} [\nu_{n,k}^T - (\nu_{n,k}^{\text{R}}(m))^+] \quad (3.35)$$

As we need to find optimal S and m for ν_n^T , we apply $\max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m}$ to the last term and $\max_{0 \leq m \leq P_n - 1}$ to the entire expression. Since $\nu_{n,k}^{\text{R}}(m) = \max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m)$ and $\nu_{n,k}^T = \max_{\ell \in \mathcal{L}_{n,k}} w_\ell$ at a leaf, adding the backhaul weight $w_{b_n}^{\text{DL}}(m)$ gives (3.32). $\square$

Lemma 4. Assume that P_n^1 serves the backhaul link of n and $\text{pr}(n)$, and receives packets from n (i.e., n transmits on b_n , BH-UL). The weight $\bar{w}_{b_n}(\psi)$ depends on $\psi = m_{\text{pr}(n)}$, which is determined by $\text{pr}(n)$'s own panel allocation. Therefore we exclude the BH-UL weight from ν_n^{R}

and let $\text{pr}(n)$ add it: at $\text{pr}(n)$'s side, the child-BH link b_n in RX mode contributes $w_{b_n}^{\text{UL}}(\psi) + \nu_n^{\text{R}}$.

$$\nu_n^{\text{R}} = \max_{1 \leq m \leq P_n} \left[\sum_{k \neq P_n^1} \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m-1} \sum_{k \in S} \max_{\ell \in \mathcal{L}_{n,k}} w_\ell - \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ \right] \quad (3.36)$$

Note that if $m = P_n$, all panels are TX mode and SI does not appear in this situation, however, for simplicity of notation, we include $m = P_n$ of max since the inner expression at $m = P_n$ reduces exactly to the correct value, $\sum_{k \neq P_n^1} \nu_{n,k}^{\text{T}}$.

Proof. Since P_n^1 is transmitting on b_n , the total TX panel count ranges from $m = 1$ to $m = P_n$. The derivation is analogous to (3.32): partition $\mathcal{P}_n \setminus P_n^1$ into $m - 1$ TX panels and the rest non-TX, then optimize over the partition and m . When $m = P_n$, all panels are TX and the $\bar{w}_\ell(m)$ terms cancel out, leaving $\sum_{k \neq P_n^1} \nu_{n,k}^{\text{T}}$. $\square$

Lemma 5. Given b_n is idle, P_n^1 contributes $\nu_{n,P_n^1}^{\text{IDLE}} = 0$ and the remaining panels are free.

$$\nu_n^{\text{IDLE}} = \max_{0 \leq m \leq P_n} \left[\sum_{k \neq P_n^1} \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m} \sum_{k \in S} \max_{\ell \in \mathcal{L}_{n,k}} w_\ell - \left(\max_{\ell \in \mathcal{L}_{n,k}} \bar{w}_\ell(m) \right)^+ \right]. \quad (3.37)$$

3.2.2 Intermediate Node

Let n be a node with at least one child IAB node. In other words, $\mathcal{L}_n$ contains at least one child-BH links.

Definition 8. (Link level) For each $c \in \mathcal{C}_n$,

$$\nu_{n,b_c}^{\text{T}} = \nu_c^{\text{T}}, \quad \nu_{n,b_c}^{\text{R}}(m) = w_{b_c}^{\text{UL}}(m) + \nu_c^{\text{R}}, \quad \nu_{n,b_c}^{\text{IDLE}} = \nu_c^{\text{IDLE}}. \quad (3.38)$$

The key differences from leaves are $w_{b_c}^{\text{UL}}(m)$ is added to ν_c^{R} and $\nu_{n,b_c}^{\text{IDLE}} = \nu_c^{\text{IDLE}}$ can be non-zero. Note that if P_c^1 selects AC link and turn off the b_c , the case is contained in $\nu_{n,b_c}^{\text{IDLE}}$.

Definition 9. (Panel level) When panel k is active, it serves exactly one link; the remaining links on k are idle and contribute $\nu_{n,\ell}^{\text{IDLE}}$ each. Hence:

$$\nu_{n,k}^{\text{IDLE}} = \sum_{\ell \in \mathcal{L}_{n,k}} \nu_{n,\ell}^{\text{IDLE}}, \quad (3.39)$$

$$\nu_{n,k}^{\text{T}} = \nu_{n,k}^{\text{IDLE}} + \max_{\ell \in \mathcal{L}_{n,k}} (\nu_{n,\ell}^{\text{T}} - \nu_{n,\ell}^{\text{IDLE}}), \quad (3.40)$$

$$\nu_{n,k}^{\text{R}}(m) = \nu_{n,k}^{\text{IDLE}} + \max_{\ell \in \mathcal{L}_{n,k}} (\nu_{n,\ell}^{\text{R}}(m) - \nu_{n,\ell}^{\text{IDLE}}). \quad (3.41)$$

At a leaf, $\nu_{n,\ell}^{\text{IDLE}} = 0$ for all ℓ , so these reduce to $\nu_{n,k}^{\text{T}} = \max_{\ell} \nu_{n,\ell}^{\text{T}}$ and $\nu_{n,k}^{\text{R}}(m) = \max_{\ell} \nu_{n,\ell}^{\text{R}}(m)$. At an intermediate node, $\nu_{n,b_c}^{\text{IDLE}} = \nu_c^{\text{IDLE}}$ can be nonzero, so $\nu_{n,k}^{\text{IDLE}}$ and the differences $\nu_{n,\ell}^{\text{T}} - \nu_{n,\ell}^{\text{IDLE}}$ must be used to correctly account for the idle contributions of other links on the same panel.

Definition 10. (Node level) For IAB n , if the parent of n , $\text{pr}(n)$ transmits through b_n to n , define ν_n^{T} as the optimal weighted sum of $\mathcal{T}_n$. Analogously, when the $\text{pr}(n)$ receives through b_n from n and b_n is idle, define ν_n^{R} and ν_n^{IDLE} as the optimal weighted sum of $\mathcal{T}_n$, respectively. Their derivations are analogous to the leaf case. The only difference is that a non-TX panel's

value is now $\max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}})$ instead of $(\nu_{n,k}^R(m))^+$, since $\nu_{n,k}^{\text{IDLE}}$ can be nonzero.

$$\begin{aligned} \nu_n^T = \max_{0 \leq m \leq P_n - 1} & \left[w_{b_n}^{\text{DL}}(m) + \nu_{n,P_n^1}^{\text{IDLE}} + \sum_{k \neq P_n^1} \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right. \\ & \left. + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m} \sum_{k \in S} \nu_{n,k}^T - \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right], \end{aligned} \quad (3.42)$$

$$\begin{aligned} \nu_n^R = \max_{1 \leq m \leq P_n} & \left[\nu_{n,P_n^1}^{\text{IDLE}} + \sum_{k \neq P_n^1} \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right. \\ & \left. + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m-1} \sum_{k \in S} \nu_{n,k}^T - \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right], \end{aligned} \quad (3.43)$$

$$\begin{aligned} \nu_n^{\text{IDLE}} = \max_{0 \leq m \leq P_n} & \left[\nu_{n,P_n^1}^{\text{IDLE}} + \sum_{k \neq P_n^1} \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right. \\ & \left. + \max_{S \subseteq \mathcal{P}_n \setminus P_n^1, |S|=m} \sum_{k \in S} \nu_{n,k}^T - \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}}) \right]. \end{aligned} \quad (3.44)$$

3.2.3 Root

Since the root r has no parent-backhaul link, only ν_r^{IDLE} is needed. There is no b_r to exclude, so all panels participate:

$$\nu_r^{\text{IDLE}} = \max_{0 \leq m \leq P_r} \left[\sum_{k \in \mathcal{P}_r} \max(\nu_{r,k}^R(m), \nu_{r,k}^{\text{IDLE}}) + \max_{S \subseteq \mathcal{P}_r, |S|=m} \sum_{k \in S} \nu_{r,k}^T - \max(\nu_{r,k}^R(m), \nu_{r,k}^{\text{IDLE}}) \right]. \quad (3.45)$$

3.2.4 Algorithm

Algorithm 1 computes the three values $(\nu_n^T, \nu_n^R, \nu_n^{\text{IDLE}})$ at a single node given the messages from its children. Algorithm 2 calls it in a forward pass and then recovers the schedule in a backward pass.

Algorithm 1 Solver for Sub-Problem at node n

Require: node n , children's messages $\{(\nu_c^T, \nu_c^R, \nu_c^{\text{IDLE}})\}_{c \in \mathcal{C}(n)}$

Ensure: $(\nu_n^T, \nu_n^R, \nu_n^{\text{IDLE}})$ and stored assignments $(m_\theta^*, S_\theta^*, \{\ell_k^T, \ell_k^R\}_k)$ for each mode θ

- 1: **Link level:** For each $\ell \in \mathcal{L}_n$, set
 - 2: Access: $\nu_{n,\ell}^T = w_\ell$, $\nu_{n,\ell}^R(m) = \bar{w}_\ell(m)$, $\nu_{n,\ell}^{\text{IDLE}} = 0$.
 - 3: Child-BH b_c : $\nu_{n,b_c}^T = \nu_c^T$, $\nu_{n,b_c}^R(m) = w_{b_c}^{\text{UL}}(m) + \nu_c^R$, $\nu_{n,b_c}^{\text{IDLE}} = \nu_c^{\text{IDLE}}$.
 - 4: Parent-BH b_n : $\nu_{n,b_n}^T = 0$, $\nu_{n,b_n}^R(m) = w_{b_n}^{\text{DL}}(m)$, $\nu_{n,b_n}^{\text{IDLE}} = 0$.
 - 5: **Panel level:** For each $k \in \mathcal{P}_n$, set
 - 6: $\nu_{n,k}^{\text{IDLE}} = \sum_{\ell \in \mathcal{L}_{n,k}} \nu_{n,\ell}^{\text{IDLE}}$.
 - 7: $\ell_k^T = \arg \max_{\ell \in \mathcal{L}_{n,k}} (\nu_{n,\ell}^T - \nu_{n,\ell}^{\text{IDLE}})$.
 - 8: $\nu_{n,k}^T = \nu_{n,k}^{\text{IDLE}} + (\nu_{n,\ell_k^T}^T - \nu_{n,\ell_k^T}^{\text{IDLE}})$.
 - 9: For each $m \in \{0, \dots, P_n\}$:
 - 10: $\ell_k^R(m) = \arg \max_{\ell \in \mathcal{L}_{n,k}} (\nu_{n,\ell}^R(m) - \nu_{n,\ell}^{\text{IDLE}})$.
 - 11: $\nu_{n,k}^R(m) = \nu_{n,k}^{\text{IDLE}} + (\nu_{n,\ell_k^R(m)}^R(m) - \nu_{n,\ell_k^R(m)}^{\text{IDLE}})$.
 - 12: **Node level:** Let $\mathcal{Q} = \mathcal{P}_n \setminus P_n^1$ (or $\mathcal{P}_n$ if $n = r$).
 - 13: **for** each mode $\theta \in \{T, R, \phi\}$ (only ϕ if $n = r$) **do**
 - 14: **for** each m in range $([0, P_n - 1]$ for T ; $[1, P_n]$ for R ; $[0, P_n]$ for ϕ) **do**
 - 15: For each $k \in \mathcal{Q}$: $d_k(m) = \nu_{n,k}^T - \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}})$.
 - 16: Sort $\{d_k(m)\}_{k \in \mathcal{Q}}$ in decreasing order.
 - 17: $m' = m$ if $\theta \in \{T, \phi\}$; $m' = m - 1$ if $\theta = R$.
 - 18: Pick top- m' panels into $S^*(m)$. {exactly m' , even if $d_k < 0$ }
 - 19: $V(\theta, m) = \sum_{k \in S^*} \nu_{n,k}^T + \sum_{k \in \mathcal{Q} \setminus S^*} \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}})$.
 - 20: **if** $n \neq r$ and $\theta = T$ **then**
 - 21: $V(\theta, m) += w_{b_n}^{\text{DL}}(m) + \nu_{n,P_n^1}^{\text{IDLE}}$.
 - 22: **else if** $n \neq r$ **then**
 - 23: $V(\theta, m) += \nu_{n,P_n^1}^{\text{IDLE}}$.
 - 24: **end if**
 - 25: **end for**
 - 26: $m_\theta^* = \arg \max_m V(\theta, m)$.
 - 27: $\nu_n^\theta = V(\theta, m_\theta^*)$.
 - 28: Store the optimal TX count m_θ^* , the TX panel set $S^*(m_\theta^*)$, and the served links $\ell_k^T, \ell_k^R(m_\theta^*)$ for each panel k .
 - 29: **end for**
 - 30: **return** $(\nu_n^T, \nu_n^R, \nu_n^{\text{IDLE}})$ and the stored assignments for each mode θ .
-

Algorithm 2 DP for Maximum Weighted Schedule

Require: $\{w_\ell\}_{\ell \in \mathcal{L}}$ **Ensure:** s^*

```
1: Phase 1 – Forward pass, from the leaves to the root.
2: for each leaf node  $n$  (all downlinks are access links) do
3:   Run Algorithm 1 with  $\mathcal{C}(n) = \emptyset$ .
4:   Send  $(\nu_n^T, \nu_n^R, \nu_n^{\text{IDLE}})$  to parent  $\text{pr}(n)$ .
5: end for
6: for each non-leaf node  $n$  of the tree, from leaves toward root do
7:   After receiving messages from all children in  $\mathcal{C}(n)$ , run Algorithm 1.
8:   Send  $(\nu_n^T, \nu_n^R, \nu_n^{\text{IDLE}})$  to parent  $\text{pr}(n)$ .
9: end for
10:  $v^* = \nu_r^{\text{IDLE}}$ .
11: Phase 2 – Backward pass, from the root to the leaves.
12:  $\theta_r \leftarrow \phi$ .  $s_\ell^* = 0, \forall \ell \in \mathcal{L}$ .
13: for  $n = 1, 2, \dots, N$  do
14:   Retrieve the stored assignments for mode  $\theta_n$ : optimal TX count  $m^*$ , TX panel set  $S^*$ , and served links  $\{\ell_k^T, \ell_k^R\}_k$ .
15:   if  $n \neq r$  and  $\theta_n = T$  then
16:      $s^*[b_n] = \text{RX}$ .
17:   else if  $n \neq r$  and  $\theta_n = R$  then
18:      $s^*[b_n] = \text{TX}$ .
19:   end if
20:   for each  $k \in \mathcal{Q}$  do
21:     if  $k \in S^*$  then
22:        $s^*[\ell_k^T] = \text{TX}$ .
23:     else if  $\nu_{n,k}^R(m^*) > \nu_{n,k}^{\text{IDLE}}$  then
24:        $s^*[\ell_k^R(m^*)] = \text{RX}$ .
25:     end if
26:   end for
27:   for each child IAB  $c$  of  $n$  do
28:     if  $k(c) \in S^*$  and  $\ell_{k(c)}^T = b_c$  then
29:        $\theta_c \leftarrow T$ .
30:     else if  $k(c) \notin S^*$  and  $\ell_{k(c)}^R = b_c$  and  $\nu_{n,k(c)}^R > \nu_{n,k(c)}^{\text{IDLE}}$  then
31:        $\theta_c \leftarrow R$ .
32:     else
33:        $\theta_c \leftarrow \phi$ .
34:     end if
35:     Send  $\theta_c$  to child  $c$ .
36:   end for
37: end for
```

3.2.5 Computational Complexity

We analyze the cost of Algorithm 1 at a single node n and then sum over the tree.

At the link level, $\nu_{n,\ell}^T$ and $\nu_{n,\ell}^{\text{IDLE}}$ are obtained in $\mathcal{O}(1)$ per link, and $\nu_{n,\ell}^R(m)$ must be evaluated for each of the $P_n + 1$ candidate values of m . Over all $|\mathcal{L}_n|$ links incident to n , this costs $\mathcal{O}(P_n \cdot |\mathcal{L}_n|)$. At the panel level, the RX value $\nu_{n,k}^R(m)$ requires an $\arg \max$ over

$|\mathcal{L}_{n,k}|$ links for each m , giving $\mathcal{O}(P_n \cdot |\mathcal{L}_{n,k}|)$ per panel. Summing over all P_n panels yields $\mathcal{O}(P_n \cdot |\mathcal{L}_n|)$, the same order as the link level. At the node level, for a fixed mode θ and a fixed m , computing $d_k(m) = \nu_{n,k}^T - \max(\nu_{n,k}^R(m), \nu_{n,k}^{\text{IDLE}})$ for every panel in $\mathcal{Q}$ takes $\mathcal{O}(P_n)$. Sorting these scores takes $\mathcal{O}(P_n \log P_n)$. After sorting, selecting the top- m' panels and evaluating $V(\theta, m)$ is $\mathcal{O}(P_n)$, so the sorting dominates. Since m ranges over at most $P_n + 1$ values, the cost per mode is $\mathcal{O}(P_n \cdot P_n \log P_n) = \mathcal{O}(P_n^2 \log P_n)$, and with three modes the node-level cost is $\mathcal{O}(P_n^2 \log P_n)$. Combining the three levels, Algorithm 1 runs in $\mathcal{O}(P_n \cdot |\mathcal{L}_n| + P_n^2 \log P_n)$ at node n . Summing over all N nodes gives the total cost of Phase 1. Each access link belongs to exactly one node's link set $\mathcal{L}_n$, while each backhaul link b_c appears in $\mathcal{L}_n$ of both its endpoints, so $\sum_n |\mathcal{L}_n| \leq 2|\mathcal{L}|$. Assuming $P_n \equiv P$ for all nodes, Phase 1 costs $\mathcal{O}(P \cdot |\mathcal{L}| + NP^2 \log P)$. Phase 2 visits each node once, reads the stored assignment in $\mathcal{O}(P_n)$, and propagates modes to children in $\mathcal{O}(|\mathcal{C}(n)|)$. Since $\sum_n |\mathcal{C}(n)| = N - 1$, the total cost of Phase 2 is $\mathcal{O}(NP)$, which is dominated by Phase 1. The overall complexity is therefore $\mathcal{O}(P \cdot |\mathcal{L}| + NP^2 \log P)$. Since P is at most 3 or 4, the overall complexity is $\mathcal{O}(|\mathcal{L}|)$.

3.3 Local Max-Weight Scheduling

The global policies of Sections 3.1.8–3.1.9 require a forward–backward message pass across the entire tree in every slot. We now develop a local policy in which each node decides its panel assignment using only its own queue lengths, the channel state on its incident links, and one scalar received from its parent.

3.3.1 Estimated Backhaul Weight

The local policy is preceded by a one-time *estimation phase*, in which the network is scheduled by the centralized full-duplex back-pressure policy of Section 3.1.8 for T_{est} slots. During this phase each parent records, for every child c , the realized BH-DL rate $R_{b_c}(\mathbf{s}, t)$ in each slot where back-pressure activates b_c . These samples calibrate the scalar weights that the local

policy later uses in place of the rates it cannot observe. The estimate is then used throughout the *operating phase*, during which every node runs the local policy below.

When a parent node n considers transmitting on BH-DL b_c to a child c , the BH-DL rate R_{b_c} depends on the child's TX count m_c , which n cannot observe. We replace this unknown rate with the fixed estimated weight

$$\hat{\mu}_{b_c} := \mathbb{E}\left[R_{b_c}(\mathbf{s}, t) \mid s_{b_c}^{\text{bp}}(t) = 1\right], \quad (3.46)$$

the conditional mean BH-DL rate when back-pressure activates b_c , computed as the empirical average over the estimation-phase samples above. All other link types have rates the scheduling node can compute itself, since for BH-UL the receiver n knows its own TX count m_n and for access links the UE endpoint causes no SI. The only residual inaccuracy of $\hat{\mu}_{b_c}$ is the finite-sample estimation error ϵ_{train} , the deviation of the empirical mean from the true conditional mean, which vanishes as the estimation phase lengthens ($T_{\text{est}} \rightarrow \infty$).

3.3.2 Policy

Definition 11. At the start of each slot, node n exchanges backhaul queue lengths with each child $c \in \mathcal{C}(n)$, sending the downlink backlog $Q_{b_c}(t)$ and receiving the uplink backlog, and observes the decision of its parent on b_n . The incident links $\mathcal{L}_n$ carry both directions, and node n may serve uplink and downlink at once across its P_n panels. Partition $\mathcal{L}_n$ into the receive links $\mathcal{R}_n$, which hold the incoming BH-DL b_n , the uplink access links, and the BH-UL from the children, and the transmit links $\mathcal{T}_n$, which hold the downlink access links, the BH-UL to the parent, and the outgoing BH-DL set $\mathcal{B}_n := \{b_c : c \in \mathcal{C}(n)\}$. Let $s_\ell(t) \in \{0, 1\}$ indicate that a panel serves link ℓ , let $m := \sum_{\ell \in \mathcal{T}_n} s_\ell$ be the transmit count, and assign each link the

weight

$$w_\ell(m) := \begin{cases} R_\ell(m, \boldsymbol{\mu}_n(t)) Q_\ell(t), & \ell \in \mathcal{R}_n, \\ R_\ell(\boldsymbol{\mu}_n(t)) Q_\ell(t), & \ell \in \mathcal{T}_n \setminus \mathcal{B}_n, \\ \hat{\mu}_{b_c} Q_{b_c}(t), & \ell = b_c \in \mathcal{B}_n. \end{cases} \quad (3.47)$$

Node n rates every receive link and every transmit link it controls exactly, while each outgoing BH-DL b_c takes the frozen estimate $\hat{\mu}_{b_c}$ in place of a rate the unobservable child transmit count would set. For each child this means node n rates the uplink it receives exactly and estimates only the downlink it transmits. The self-interference of the m transmit panels degrades every receive rate $R_\ell(m, \boldsymbol{\mu}_n)$. When the parent activates b_n (Mode T), node n solves

$$\begin{aligned} \max_{s, m} \quad & \sum_{\ell \in \mathcal{L}_n} w_\ell(m, \mu_\ell) s_\ell \\ \text{s.t.} \quad & s_{b_n} = 1, \\ & \sum_{\ell \in \mathcal{L}_n} s_\ell \leq P_n, \\ & m = \sum_{\ell \in \mathcal{T}_n} s_\ell, \\ & s_\ell \in \{0, 1\}, \quad \ell \in \mathcal{L}_n. \end{aligned} \quad (3.48)$$

The constraint $s_{b_n} = 1$ forces one panel to receive on b_n , and the budget $\sum_\ell s_\ell \leq P_n$ serves at most one link per panel. When the parent is idle (Mode ϕ), node n solves the same program with b_n removed from $\mathcal{L}_n$ and the constraint $s_{b_n} = 1$ dropped, so all P_n panels are free. The transmit count m enters the receive weights, so node n solves (3.48) by sweeping m with the sub-solver of Algorithm 1.

Why the target load exceeds the operating load. During the operating phase, the parent never sees the true BH-DL rate, since that rate depends on the child's panel assignment and self-interference, which the parent cannot observe. It instead decides whether to activate BH using the single calibrated weight from the estimation phase. Because this weight is a

fixed average rather than the rate realized in the current slot, the parent's on-off decisions are occasionally wrong. Sometimes it activates the backhaul when serving access would have carried more traffic, and sometimes it skips a backhaul opportunity that was in fact worth taking. Each such decision spends panel-time on the less useful link, so a fraction of the network's capacity is lost to acting on partial information rather than the exact state. To keep every queue stable in spite of this loss, the network is provisioned with headroom, that is, it is designed to carry a load below the rate its links could sustain if every backhaul decision were exact. The size of this headroom reflects the cost of partial information. It shrinks as the estimation phase lengthens and the calibrated weight tracks the true rate more closely, and as the self-interference penalty weakens so that the gap between activating and skipping the backhaul becomes less sensitive to the child's behavior.

Algorithm 3 FD-LocalMW-BP at node n (non-root)

Require: Local queues $\{Q_\ell\}_{\ell \in \mathcal{L}_n}$, channel state $\boldsymbol{\mu}_n(t)$, parent's BH queue Q_{b_n} , parent's BH schedule s_{b_n} , estimated weights $\{\hat{\mu}_{b_c}\}_{c \in \mathcal{C}(n)}$

Ensure: Panel assignment $\mathbf{s}_n(t)$

- 1: **TX-side link weight:** $w_\ell^{\text{TX}} \leftarrow R_\ell(t) Q_\ell$ for access links, and $w_{b_c}^{\text{TX}} \leftarrow \hat{\mu}_{b_c} Q_{b_c}$ for each outgoing BH-DL b_c .
- 2: **Send** $Q_{b_c}(t)$ to each child $c \in \mathcal{C}(n)$.
- 3: **Observe** parent's decision on b_n .
- 4: **if** parent TX on b_n (**Mode T**) **then**
- 5: Pin $P_n^1 \leftarrow \text{RX on } b_n$.
- 6: **for** $m = 0$ **to** $P_n - 1$ **do**
- 7: $V_{\text{BH}}(m) \leftarrow R_{b_n}(m, \boldsymbol{\mu}_n(t)) \cdot Q_{b_n}$. {exact incoming BH-DL rate at SI level m }
- 8: **for** each $k \in \mathcal{P}_n \setminus P_n^1$ **do**
- 9: $\nu_{n,k}^{\text{T}} \leftarrow \max_{\ell \in \mathcal{L}_{n,k}^{\text{TX}}} w_\ell^{\text{TX}}$. {access at exact rate, outgoing BH-DL at estimate}
- 10: $\nu_{n,k}^{\text{R}}(m) \leftarrow \max_{\ell \in \mathcal{L}_{n,k}^{\text{RX}}} R_\ell(m, t) \cdot Q_\ell$.
- 11: **end for**
- 12: Sort panels by marginal gain $\nu_{n,k}^{\text{T}} - (\nu_{n,k}^{\text{R}}(m))^+$, pick top m .
- 13: $V_{\text{AC}}(m) \leftarrow \sum_k (\nu_{n,k}^{\text{R}}(m))^+ + \text{top-}m \text{ gains}$.
- 14: $V(m) \leftarrow V_{\text{BH}}(m) + V_{\text{AC}}(m)$.
- 15: **end for**
- 16: $m^* \leftarrow \arg \max_m V(m)$.
- 17: Assign top- m^* gain panels to TX, the rest to their best RX link.
- 18: **else** {**Mode** ϕ : parent idle, all panels free}
- 19: **for** each $\ell \in \mathcal{L}_n$ **do**
- 20: **if** $\ell = b_c$ is an outgoing BH-DL **then**
- 21: $w_\ell \leftarrow \hat{\mu}_{b_c} \cdot Q_{b_c}$. {estimated weight: child's SI unknown to n }
- 22: **else**
- 23: $w_\ell \leftarrow R_\ell(t) \cdot Q_\ell$. {exact rate}
- 24: **end if**
- 25: **end for**
- 26: Run Algorithm 1 with weights $\{w_\ell\}$ to assign panels.
- 27: **end if**

3.3.3 Stability Region of Full-Duplex Local Scheduling

In this theoretical analysis on the stability region, we will consider only the downlink for simplicity. The analysis of the uplink can be analogously investigated, and the framework can be easily extended to combined operations.

Definition 12. For each node n , Mode T sub-modes are indexed by TX count $m \in \{0, \dots, P_n - 1\}$. In Mode T - m , panel P_n^1 is pinned to BH-DL RX (rate $R_{b_n}(m, \boldsymbol{\mu}_n)$), and m

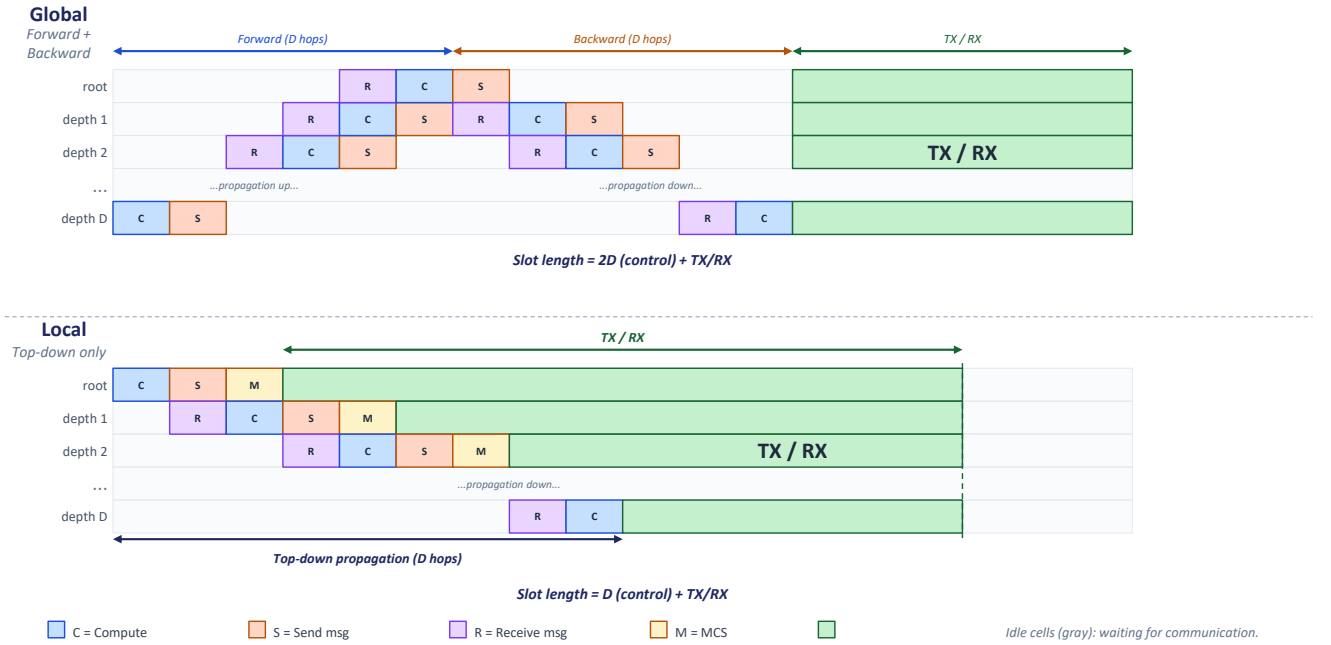


Figure 3.1: Timing diagram of global versus local scheduling on a tree of depth D . Global scheduling requires a forward pass and a backward pass before TX/RX can begin, consuming $2D$ control intervals per slot. Local scheduling propagates only top-down consuming D control intervals per slot.

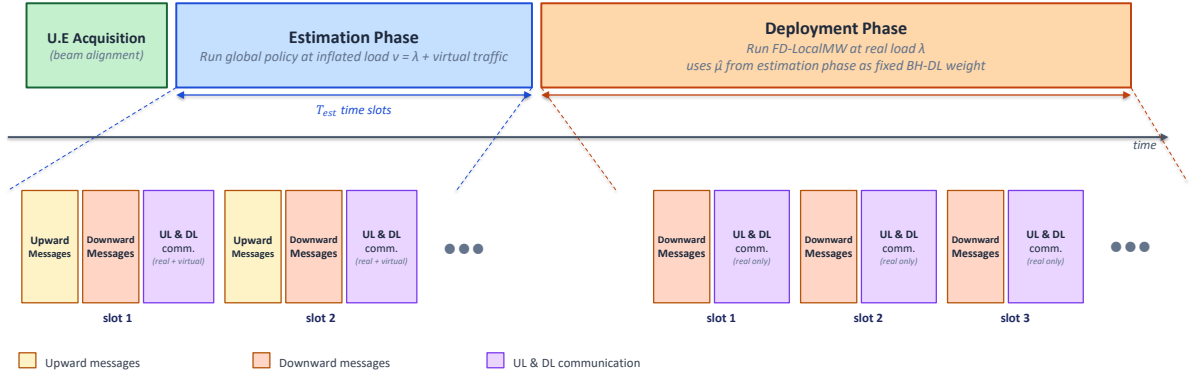


Figure 3.2: Operational timeline of FD-LocalMW-BP. After UE acquisition, the network runs the global policy at the inflated load $\lambda = \eta +$ virtual traffic for T_{est} slots (estimation phase) to produce the BH-DL weight estimates $\{\hat{\mu}_{b_c}\}$. The network then switches to FD-LocalMW-BP at the real load η (deployment phase), using $\hat{\mu}$ as the fixed BH-DL weight.

of the remaining panels are TX, creating SI level m . Define:

$$\mathcal{S}_n^{\text{T},m} := \{\mathbf{s}_n \in \{0,1\}^{|\mathcal{L}_n|} : \sum_{\ell \in \mathcal{L}_{n,k}} s_\ell \leq 1 \ \forall \text{ free panel } k, \ \sum_{\ell \in \mathcal{L}_n^{\text{TX}}} s_\ell = m\}, \quad (3.49)$$

$$\mathcal{S}_n^{\text{IDLE}} := \{\mathbf{s}_n \in \{0,1\}^{|\mathcal{L}_n|} : \sum_{\ell \in \mathcal{L}_{n,k}} s_\ell \leq 1 \ \forall \text{ panel } k\}, \quad (3.50)$$

where $\mathcal{L}_n = \mathcal{L}_n^{\text{AC}} \cup \{b_c\}_{c \in \mathcal{C}(n)}$ includes access links and child-BH links, and $\mathcal{L}_{n,k}$ is the set of links hostable on panel k .

Definition 13. For each non-root node n , define the *mode set*

$$\mathcal{M}_n := \{0, 1, \dots, P_n - 1\} \cup \{\phi\}. \quad (3.51)$$

A mode $\theta_n = m \in \{0, \dots, P_n - 1\}$ means node n receives BH-DL (panel P_n^1 pinned to RX) with m TX panels active, creating SI level m . The mode $\theta_n = \phi$ means no BH-DL reception (all panels free). The root is always in mode ϕ : $\mathcal{M}_r := \{\phi\}$. The per-mode feasible schedule set is written uniformly as $\mathcal{S}_n(\theta_n)$:

$$\mathcal{S}_n(\theta_n) := \begin{cases} \mathcal{S}_n^{\text{T},m} & \text{if } \theta_n = m \in \{0, \dots, P_n - 1\}, \\ \mathcal{S}_n^{\text{IDLE}} & \text{if } \theta_n = \phi, \end{cases} \quad (3.52)$$

where $\mathcal{S}_n^{\text{T},m}$ and $\mathcal{S}_n^{\text{IDLE}}$ are as in (3.49)–(3.50). The BH-DL rate in mode θ_n is

$$R_{b_n}(\theta_n, \boldsymbol{\mu}_n) := \begin{cases} R_{b_n}(m, \boldsymbol{\mu}_n) & \text{if } \theta_n = m \in \{0, \dots, P_n - 1\}, \\ 0 & \text{if } \theta_n = \phi. \end{cases} \quad (3.53)$$

Definition 14. A *feasible network schedule* is a pair $\sigma = (\boldsymbol{\theta}, \mathbf{s})$ consisting of:

- (i) a mode assignment $\boldsymbol{\theta} = (\theta_n)_{n \in \mathcal{B}}$ with $\theta_n \in \mathcal{M}_n$ for each node,
- (ii) a link activation $\mathbf{s} = (\mathbf{s}_n)_{n \in \mathcal{B}}$ with $\mathbf{s}_n \in \{0,1\}^{|\mathcal{L}_n|}$ for each node,

subject to:

(C1) *Feasibility*: $\mathbf{s}_n \in \mathcal{S}_n(\theta_n)$ for every $n \in \mathcal{B}$.

(C2) *Parent-child coupling*: for each parent-child pair (n, c) ,

$$\theta_c \neq \phi \iff s_{b_c} = 1 \text{ at node } n. \quad (3.54)$$

Let Σ_{net} denote the set of all such pairs. It is finite.

Each $\sigma = (\boldsymbol{\theta}, \mathbf{s}) \in \Sigma_{\text{net}}$ and channel state $\boldsymbol{\mu} = (\mu_\ell)_{\ell \in \mathcal{L}}$ jointly produce a *rate vector* $\mathbf{r}(\sigma, \boldsymbol{\mu}) \in \mathbb{R}_{\geq 0}^{|\mathcal{A}|+|\mathcal{B}_{\text{BH}}|}$:

$$r_{b_n}(\sigma, \boldsymbol{\mu}) = R_{b_n}(\theta_n, \boldsymbol{\mu}_n), \quad r_\ell(\sigma, \boldsymbol{\mu}) = \mu_\ell \cdot s_\ell^\sigma, \quad \forall b_n \in \mathcal{B}_{\text{BH}}, \ell \in \mathcal{A}. \quad (3.55)$$

Let $\mathcal{M}$ denote the set of all possible channel states and $\pi_\boldsymbol{\mu} := \Pr[\boldsymbol{\mu}(t) = \boldsymbol{\mu}]$ the stationary distribution.

Definition 15. A scheduling policy π is *local* if each node's scheduling decision depends only on its own queue state, its own channel state, and the mode signal received from its parent. Formally, for each parent-child pair (p, n) :

$$\mathbb{P}^\pi(s_{b_n}(t) = 1 \mid \boldsymbol{\mu}_n(t), \mathbf{Q}(t)) = \mathbb{P}^\pi(s_{b_n}(t) = 1 \mid \mathbf{Q}(t)), \quad (3.56)$$

i.e., parent p 's BH-DL activation for child n is independent of child n 's channel state $\boldsymbol{\mu}_n(t)$. Within a mode, each node may observe its own channel $\boldsymbol{\mu}_n(t)$ and adapt its panel assignment accordingly.

Definition 16. The *global stability region* is the set of arrival rate vectors achievable by some policy:

$$\Lambda := \{\boldsymbol{\lambda} \geq 0 : \exists \text{ a policy } \pi \text{ with } \bar{S}_\ell^\pi \geq \lambda_\ell \forall \ell\}. \quad (3.57)$$

The *local stability region* is the set achievable by local policies:

$$\Lambda^{\text{loc}} := \{\boldsymbol{\lambda} \geq 0 : \exists \text{ a local policy } \pi \text{ with } \bar{S}_\ell^\pi \geq \lambda_\ell \forall \ell\}. \quad (3.58)$$

Here $\bar{S}_\ell^\pi := \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}^\pi[\mu_\ell(t) s_\ell(t)]$ is the time-averaged service rate on link ℓ under π .

Theorem 3. (Carathéodory [16, Theorem 17.1].) *Let $S \subset \mathbb{R}^d$ be a finite set. If $\mathbf{x} \in \text{Conv}(S)$, then there exist points $\mathbf{s}_1, \dots, \mathbf{s}_k \in S$ with $k \leq d + 1$ and weights $p_1, \dots, p_k \geq 0$ with $\sum_{j=1}^k p_j = 1$ such that $\mathbf{x} = \sum_{j=1}^k p_j \mathbf{s}_j$.*

Lemma 6. *Let π be any scheduling policy, and let $E_1, \dots, E_K$ be finitely many events each determined by $(\mathbf{Q}(t), \boldsymbol{\mu}(t), \mathbf{s}(t))$. There exists a subsequence $\{T_k\}_{k \geq 1}$ with $T_k \rightarrow \infty$ such that*

$$\lim_{k \rightarrow \infty} \frac{1}{T_k} \sum_{t=0}^{T_k-1} \mathbb{P}^\pi(E_i(t)) \quad (3.59)$$

exists for every $i = 1, \dots, K$.

Proof. Each Cesàro mean $\bar{p}_T(E_i) := \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{P}^\pi(E_i(t))$ lies in $[0, 1]$, so the sequence $\{\bar{p}_T(E_i)\}_{T \geq 1}$ is bounded for each i . By the Bolzano–Weierstrass theorem, $\{\bar{p}_T(E_1)\}$ has a convergent subsequence $\{T_k^{(1)}\}$. Restricting to this subsequence, $\{\bar{p}_{T_k^{(1)}}(E_2)\}$ is again bounded and has a further convergent subsequence $\{T_k^{(2)}\}$. Proceeding by diagonal extraction over $i = 1, \dots, K$ yields a single subsequence $\{T_k\}$ along which all K limits exist simultaneously. $\square$

Definition 17. For a stable local policy π , define the total Mode T fraction $\rho_n := \lim_{T_k \rightarrow \infty} T_k^{-1} \sum_{t=0}^{T_k-1} \mathbb{P}^\pi(\theta_n(t) \neq \phi)$, and for each mode m , channel state $\boldsymbol{\mu}_n$, and link ℓ :

$$\alpha_n(m, \boldsymbol{\mu}_n) := \lim_{T_k \rightarrow \infty} \frac{1}{T_k} \sum_{t=0}^{T_k-1} \mathbb{P}^\pi(\theta_n(t) = m, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n), \quad (3.60)$$

$$y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) := \lim_{T_k \rightarrow \infty} \frac{1}{T_k} \sum_{t=0}^{T_k-1} \mathbb{P}^\pi(s_\ell(t)=1, \theta_n(t) = m, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n), \quad (3.61)$$

$$y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) := \lim_{T_k \rightarrow \infty} \frac{1}{T_k} \sum_{t=0}^{T_k-1} \mathbb{P}^\pi(s_\ell(t)=1, \theta_n(t) = \phi, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n), \quad (3.62)$$

where $\{T_k\}$ is a convergent subsequence guaranteed by Lemma 6 (chosen jointly for all finitely many events via a diagonal argument). The marginal mode fraction is $\alpha_n(m) := \sum_{\boldsymbol{\mu}_n} \alpha_n(m, \boldsymbol{\mu}_n)$.

LP Characterization. For following works, we introduce Linear Programming characterization of stability region since it is intractable to handle convex hull form. Let $\rho_n \in [0, 1]$ denote the total Mode T fraction at node n (channel-independent by the locality constraint). Within Mode T , the child observes its channel $\boldsymbol{\mu}_n$ and chooses the TX count m ; we write

$$\alpha_n(m, \boldsymbol{\mu}_n) := \mathbb{P}(\theta_n = m, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n) \quad (3.63)$$

for the joint mode-channel fraction, and

$$y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) := \mathbb{P}(s_\ell=1, \theta_n=m, \boldsymbol{\mu}_n(t)=\boldsymbol{\mu}_n), \quad (3.64)$$

$$y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) := \mathbb{P}(s_\ell=1, \theta_n=\phi, \boldsymbol{\mu}_n(t)=\boldsymbol{\mu}_n) \quad (3.65)$$

for the within-mode, channel-conditional activation fractions. These are long-run time fractions obtained from any stationary local policy.

Proposition 1. $\boldsymbol{\lambda} \in \Lambda^{\text{loc}}$ if and only if the following LP in variables $(\rho_n, \alpha_n(m, \boldsymbol{\mu}_n), y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n), y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \geq$

0) is feasible:

$$\sum_m \sum_{\boldsymbol{\mu}_n} \alpha_n(m, \boldsymbol{\mu}_n) R_{b_n}(m, \boldsymbol{\mu}_n) \geq \lambda_{b_n}, \quad \forall n \neq r, \quad (\text{LP-BH})$$

$$\sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) + \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \geq \lambda_\ell, \quad \forall n, \ell \in \mathcal{L}_n^{\text{AC}}, \quad (\text{LP-AC})$$

$$\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) \leq \alpha_n(m, \boldsymbol{\mu}_n), \quad \forall n, m, \text{free } k, \boldsymbol{\mu}_n, \quad (\text{LP-PNL}^{\text{T}})$$

$$\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \leq (1 - \rho_n) \pi_{\boldsymbol{\mu}_n}, \quad \forall n, \text{panel } k, \boldsymbol{\mu}_n, \quad (\text{LP-PNL}^{\text{IDLE}})$$

$$\sum_{\ell \in \mathcal{L}_n^{\text{TX}}} y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) = m \alpha_n(m, \boldsymbol{\mu}_n), \quad \forall n, m, \boldsymbol{\mu}_n, \quad (\text{LP-TX})$$

$$\sum_m \alpha_n(m, \boldsymbol{\mu}_n) = \rho_n \pi_{\boldsymbol{\mu}_n}, \quad \forall n, \boldsymbol{\mu}_n, \quad (\text{LP-IND})$$

$$\sum_m \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{T},m}(\boldsymbol{\mu}_p) + \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{IDLE}}(\boldsymbol{\mu}_p) = \rho_c, \quad \forall (p, c), \quad (\text{LP-PC})$$

$$\rho_n \leq 1, \quad \forall n. \quad (\text{LP-}\rho)$$

The constraints encode the following physical and structural requirements:

- (LP-BH): the average BH-DL service rate, jointly over SI level m and channel state $\boldsymbol{\mu}_n$, must meet the BH arrival rate.
- (LP-AC): the average access service rate, summed over all modes and channel states, must meet the access arrival rate. Within-mode channel adaptation is captured by the $\boldsymbol{\mu}_n$ variables.
- (LP-PNL^T)–(LP-PNL^{IDLE}): each panel serves at most one link per slot. The inequality (rather than equality) allows a panel to be idle — for instance, when no link on that panel has positive weight. The total link activation on panel k cannot exceed the time fraction spent in the corresponding mode and channel state.

- (LP-TX): in Mode T - m , exactly m of the free panels are in TX mode, regardless of the channel state.
- (LP-IND): the sole locality constraint. It forces the total Mode T fraction ρ_n to be independent of the child's channel state $\boldsymbol{\mu}_n$, encoding the local policy condition (3.56). Removing (LP-IND) and allowing $\sum_m \alpha_n(m, \boldsymbol{\mu}_n)$ to vary freely with $\boldsymbol{\mu}_n$ yields the LP for Λ .
- (LP-PC): the total BH-DL activation fraction at the parent (summed over parent's modes and channel states) must equal the child's Mode T fraction ρ_c .
- (LP- ρ): the Mode T fraction cannot exceed 1.

At the root: (LP-BH) is absent and $\rho_r = 0$.

Proof. ($\Rightarrow$) Stabilizing local policy implies LP feasibility. Let ϖ be a local policy (Definition 15) that stabilizes $\boldsymbol{\lambda}$, i.e.,

$$\limsup_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}^{\varpi}[Q_{\ell}(t)] < \infty \quad \forall \ell. \quad (3.66)$$

(We write ϖ for the policy to distinguish it from the channel state probability $\pi_{\boldsymbol{\mu}_n} := \Pr[\boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n]$.) Define the occupation variables $(\rho_n, \alpha_n(m, \boldsymbol{\mu}_n), y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n), y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n))$ as the Cesàro subsequential limits of Definition 17 (which exist by Lemma 6). We verify that these variables satisfy each LP constraint:

(LP-IND): Since ϖ is local (Definition 15), the parent's BH activation decision is independent of the child's channel: $\mathbb{P}^{\varpi}(\theta_n \neq \phi \mid \boldsymbol{\mu}_n) = \mathbb{P}^{\varpi}(\theta_n \neq \phi) = \rho_n$. Hence $\sum_m \alpha_n(m, \boldsymbol{\mu}_n) = \mathbb{P}^{\varpi}(\theta_n \neq \phi, \boldsymbol{\mu}_n) = \rho_n \pi_{\boldsymbol{\mu}_n}$.

(LP-PNL^T): In every slot, each panel serves at most one link: $\sum_{\ell \in \mathcal{L}_{n,k}} s_{\ell}(t) \leq 1$. Multiply both sides by the indicator $\mathbf{1}\{\theta_n(t) = m, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n\}$:

$$\sum_{\ell \in \mathcal{L}_{n,k}} s_{\ell}(t) \mathbf{1}\{\theta_n(t)=m, \boldsymbol{\mu}_n(t)=\boldsymbol{\mu}_n\} \leq \mathbf{1}\{\theta_n(t)=m, \boldsymbol{\mu}_n(t)=\boldsymbol{\mu}_n\}. \quad (3.67)$$

Take expectations and the Cesàro mean $\frac{1}{T} \sum_{t=0}^{T-1}$:

$$\sum_{\ell \in \mathcal{L}_{n,k}} \lim_{T \rightarrow \infty} \underbrace{\frac{1}{T} \sum_{t=0}^{T-1} \mathbb{P}(s_\ell=1, \theta_n=m, \boldsymbol{\mu}_n)}_{=y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n)} \leq \underbrace{\lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{P}(\theta_n=m, \boldsymbol{\mu}_n)}_{=\alpha_n(m, \boldsymbol{\mu}_n)}. \quad (3.68)$$

(LP-PNL^{IDLE}): Identically, multiplying $\sum_\ell s_\ell \leq 1$ by $\mathbf{1}\{\theta_n=\phi, \boldsymbol{\mu}_n\}$, taking expectations and Cesàro means:

$$\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \leq \underbrace{\lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{P}(\theta_n=\phi, \boldsymbol{\mu}_n)}_{=(1-\rho_n) \pi_{\boldsymbol{\mu}_n}}. \quad (3.69)$$

(LP-TX): When $\theta_n(t) = m$, exactly m TX panels are active: $\sum_{\ell \in \mathcal{L}_n^{\text{TX}}} s_\ell(t) = m$. Multiply by $\mathbf{1}\{\theta_n=m, \boldsymbol{\mu}_n\}$ and average:

$$\sum_{\ell \in \mathcal{L}_n^{\text{TX}}} y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) = m \cdot \alpha_n(m, \boldsymbol{\mu}_n). \quad (3.70)$$

(LP-PC): Child c is in Mode T iff parent p activates b_c : $\mathbb{P}(s_{b_c} = 1) = \mathbb{P}(\theta_c \neq \phi) = \rho_c$. Decomposing the left side over parent's modes and channel states:

$$\sum_m \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{T},m}(\boldsymbol{\mu}_p) + \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{IDLE}}(\boldsymbol{\mu}_p) = \mathbb{P}(s_{b_c} = 1) = \rho_c. \quad (3.71)$$

(LP-BH): The per-slot BH-DL service is $R_{b_n}(m_n(t), \boldsymbol{\mu}_n(t)) \cdot \mathbf{1}\{\theta_n(t) \neq \phi\}$. Decomposing

over joint events and averaging:

$$\begin{aligned}
\bar{S}_{b_n} &= \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[R_{b_n}(m_n(t), \boldsymbol{\mu}_n(t)) \mathbf{1}\{\theta_n(t) \neq \phi\}] \\
&= \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \sum_m \sum_{\boldsymbol{\mu}_n} R_{b_n}(m, \boldsymbol{\mu}_n) \mathbb{P}(\theta_n(t)=m, \boldsymbol{\mu}_n(t)=\boldsymbol{\mu}_n) \\
&= \sum_m \sum_{\boldsymbol{\mu}_n} R_{b_n}(m, \boldsymbol{\mu}_n) \alpha_n(m, \boldsymbol{\mu}_n).
\end{aligned} \tag{3.72}$$

By mean stability, $\bar{S}_{b_n} \geq \lambda_{b_n}$: otherwise $\bar{S}_{b_n} < \lambda_{b_n}$, i.e., $\bar{S}_{b_n} = \lambda_{b_n} - \epsilon$ for some $\epsilon > 0$, which gives

$$\lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[\Delta q_{b_n}(t)] \geq \lambda_{b_n} - \bar{S}_{b_n} = \epsilon > 0, \tag{3.73}$$

implying $\mathbb{E}[q_{b_n}(T)]/T \geq \epsilon$, contradicting mean stability. Hence (LP-BH) holds.

(LP-AC): The access service on link ℓ is $\mu_\ell(t) \cdot s_\ell(t)$. Decomposing over joint events and averaging:

$$\begin{aligned}
\bar{S}_\ell &= \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \mathbb{E}[\mu_\ell(t) s_\ell(t)] \\
&= \lim_{T \rightarrow \infty} \frac{1}{T} \sum_{t=0}^{T-1} \sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell \mathbb{P}(s_\ell=1, \theta_n=m, \boldsymbol{\mu}_n) + \frac{1}{T} \sum_t \sum_{\boldsymbol{\mu}_n} \mu_\ell \mathbb{P}(s_\ell=1, \theta_n=\phi, \boldsymbol{\mu}_n) \\
&= \sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) + \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n).
\end{aligned} \tag{3.74}$$

By mean stability, $\bar{S}_\ell \geq \lambda_\ell$: otherwise $\bar{S}_\ell = \lambda_\ell - \epsilon$ for some $\epsilon > 0$, and $\frac{1}{T} \sum_t \mathbb{E}^\varpi[\Delta q_\ell^f(t)] \geq \epsilon > 0$ for some flow f traversing link ℓ , implying $\mathbb{E}^\varpi[q_\ell^f(T)]/T \geq \epsilon$, which contradicts mean stability. Hence (LP-AC) holds.

(LP- ρ): $\rho_n = \mathbb{P}(\theta_n \neq \phi) \leq 1$.

($\Leftarrow$) *LP feasibility implies $\boldsymbol{\lambda} \in \Lambda^{\text{loc}}$.* Given a feasible solution $(\rho_n, \alpha_n(m, \boldsymbol{\mu}_n), y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n), y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n))$, we construct an explicit stationary local policy.

Step 1. For each $n, m, \boldsymbol{\mu}_n$ with $\alpha_n(m, \boldsymbol{\mu}_n) > 0$, define

$$z_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) := \frac{y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n)}{\alpha_n(m, \boldsymbol{\mu}_n)}. \quad (3.75)$$

This is the conditional probability of activating link ℓ , given mode m and channel $\boldsymbol{\mu}_n$. From (LP-PNL^T): $\sum_{\ell \in \mathcal{L}_{n,k}} z_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) \leq 1$ for each panel k (panel exclusion). From (LP-TX): $\sum_{\ell \in \mathcal{L}_n^{\text{TX}}} z_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) = m$ (TX count). Together these imply $\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n) \in \text{Conv}(\mathcal{S}_n^{\text{T},m})$. Define $\mathbf{z}_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n)$ analogously; by (LP-PNL^{IDLE}), $\mathbf{z}_n^{\text{IDLE}}(\boldsymbol{\mu}_n) \in \text{Conv}(\mathcal{S}_n^{\text{IDLE}})$.

Step 2. Since $\mathcal{S}_n^{\text{T},m}$ is a finite set, $\text{Conv}(\mathcal{S}_n^{\text{T},m}) \subset \mathbb{R}^{|\mathcal{L}_n|}$. By Carathéodory's theorem, each $\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n) \in \text{Conv}(\mathcal{S}_n^{\text{T},m})$ can be written as a convex combination of at most $|\mathcal{L}_n| + 1$ vertices:

$$\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n) = \sum_{j=1}^J p_j^{\text{T},m}(\boldsymbol{\mu}_n) \mathbf{s}_j, \quad \mathbf{s}_j \in \mathcal{S}_n^{\text{T},m}, \quad p_j \geq 0, \quad \sum_j p_j = 1. \quad (3.76)$$

This provides a probability distribution $\{p_j^{\text{T},m}(\boldsymbol{\mu}_n)\}$ supported on at most $|\mathcal{L}_n| + 1$ schedules from $\mathcal{S}_n^{\text{T},m}$; the remaining schedules in $\mathcal{S}_n^{\text{T},m}$ receive probability zero. Similarly decompose $\mathbf{z}_n^{\text{IDLE}}(\boldsymbol{\mu}_n)$ into $\{p_j^{\text{IDLE}}(\boldsymbol{\mu}_n)\}$ over $\mathcal{S}_n^{\text{IDLE}}$.

Step 3. The LP solution and the Carathéodory decompositions from Steps 1–2 define a stationary randomized local policy $\tilde{\pi}$ as follows. For each node n , $\tilde{\pi}$ is specified by three components:

- (a) *Mode selection distribution* (for non-root n entering Mode T): given channel $\boldsymbol{\mu}_n$, select TX count m with probability

$$\tilde{\pi}_n^{\text{mode}}(m \mid \boldsymbol{\mu}_n) := \frac{\alpha_n(m, \boldsymbol{\mu}_n)}{\rho_n \pi_{\boldsymbol{\mu}_n}}, \quad (3.77)$$

which sums to 1 over m by (LP-IND);

- (b) *Mode T schedule distribution*: given mode m and channel $\boldsymbol{\mu}_n$, draw schedule $\mathbf{s}_n \in \mathcal{S}_n^{\text{T},m}$

with probability

$$\tilde{\pi}_n^{\text{T},m}(\mathbf{s}_j \mid \boldsymbol{\mu}_n) := p_j^{\text{T},m}(\boldsymbol{\mu}_n), \quad j = 1, \dots, J_{n,m,\boldsymbol{\mu}_n}, \quad (3.78)$$

where $\{p_j^{\text{T},m}(\boldsymbol{\mu}_n), \mathbf{s}_j\}$ is the Carathéodory decomposition of $\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n)$ from Step 2, with $J_{n,m,\boldsymbol{\mu}_n} \leq |\mathcal{L}_n| + 1$;

(c) *Mode ϕ schedule distribution*: given channel $\boldsymbol{\mu}_n$, draw schedule $\mathbf{s}_n \in \mathcal{S}_n^{\text{IDLE}}$ with probability

$$\tilde{\pi}_n^{\text{IDLE}}(\mathbf{s}_j \mid \boldsymbol{\mu}_n) := p_j^{\text{IDLE}}(\boldsymbol{\mu}_n), \quad j = 1, \dots, J_{n,\phi,\boldsymbol{\mu}_n}, \quad (3.79)$$

where $\{p_j^{\text{IDLE}}(\boldsymbol{\mu}_n), \mathbf{s}_j\}$ is the Carathéodory decomposition of $\mathbf{z}_n^{\text{IDLE}}(\boldsymbol{\mu}_n)$ from Step 2.

In each slot t , $\tilde{\pi}$ maps the observation $(\theta_n(t), \boldsymbol{\mu}_n(t))$ to a random schedule $\mathbf{s}_n(t)$, independently of the queue state $\mathbf{Q}(t)$ and of all other slots.

Step 4. The tree-wide schedule is constructed depth by depth:

- (i) *Root r* : observe $\boldsymbol{\mu}_r(t)$, draw $\mathbf{s}_r \sim \tilde{\pi}_r^{\text{IDLE}}(\cdot \mid \boldsymbol{\mu}_r(t))$. This determines which BH-DL links b_c are activated ($s_{b_c} = 1$).
- (ii) *Non-root n in Mode T* : observe $\boldsymbol{\mu}_n(t)$, draw $m \sim \tilde{\pi}_n^{\text{mode}}(\cdot \mid \boldsymbol{\mu}_n(t))$, then draw $\mathbf{s}_n \sim \tilde{\pi}_n^{\text{T},m}(\cdot \mid \boldsymbol{\mu}_n(t))$.
- (iii) *Non-root n in Mode ϕ* : observe $\boldsymbol{\mu}_n(t)$, draw $\mathbf{s}_n \sim \tilde{\pi}_n^{\text{IDLE}}(\cdot \mid \boldsymbol{\mu}_n(t))$.
- (iv) Repeat for deeper nodes.

Step 5. Each node uses only (a) its mode θ_n (received from parent) and (b) its own channel $\boldsymbol{\mu}_n(t)$. The total Mode T probability at child n is:

$$\mathbb{P}^{\tilde{\pi}}(\theta_n \neq \phi \mid \boldsymbol{\mu}_n) = \rho_n \quad \forall \boldsymbol{\mu}_n, \quad (3.80)$$

since the parent's BH decision does not depend on $\boldsymbol{\mu}_n$ (the parent draws from $\{p_j(\boldsymbol{\mu}_p)\}$ which is a function of $\boldsymbol{\mu}_p$, independent of $\boldsymbol{\mu}_n$). Hence $\tilde{\pi}$ is local.

Step 6. By (LP-PC), the probability that parent p activates b_c (summed over parent modes and channels) equals ρ_c , matching child c 's Mode T fraction.

Step 7. We verify that the time-averaged service rates under $\tilde{\pi}$ recover the LP values. Under $\tilde{\pi}$, the joint probability that node n is in mode m with channel $\boldsymbol{\mu}_n$ and link ℓ is active is:

$$\begin{aligned}\mathbb{P}^{\tilde{\pi}}(s_\ell=1, \theta_n=m, \boldsymbol{\mu}_n) &= \underbrace{\mathbb{P}(\theta_n=m, \boldsymbol{\mu}_n)}_{\alpha_n(m, \boldsymbol{\mu}_n)} \cdot \underbrace{\mathbb{P}^{\tilde{\pi}}(s_\ell=1 \mid \theta_n=m, \boldsymbol{\mu}_n)}_{=\sum_j \tilde{\pi}_n^{\text{T},m}(\mathbf{s}_j \mid \boldsymbol{\mu}_n) \cdot [\mathbf{s}_j]_\ell} \\ &= \alpha_n(m, \boldsymbol{\mu}_n) \cdot [\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n)]_\ell \\ &= \alpha_n(m, \boldsymbol{\mu}_n) \cdot \frac{y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n)}{\alpha_n(m, \boldsymbol{\mu}_n)} = y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n),\end{aligned}\tag{3.81}$$

where $[\mathbf{z}]_\ell$ denotes the ℓ -th component and the second equality uses $\sum_j p_j \mathbf{s}_j = \mathbf{z}$ from the Carathéodory decomposition. Similarly, $\mathbb{P}^{\tilde{\pi}}(s_\ell=1, \theta_n=\phi, \boldsymbol{\mu}_n) = y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n)$.

For BH-DL at non-root n : since $s_{b_n} = 1$ whenever $\theta_n \neq \phi$,

$$\bar{S}_{b_n} = \sum_m \sum_{\boldsymbol{\mu}_n} R_{b_n}(m, \boldsymbol{\mu}_n) \alpha_n(m, \boldsymbol{\mu}_n) \geq \lambda_{b_n} \quad (\text{by (LP-BH)}).\tag{3.82}$$

For access link $\ell \in \mathcal{L}_n^{\text{AC}}$:

$$\bar{S}_\ell = \sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) + \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \geq \lambda_\ell \quad (\text{by (LP-AC)}).\tag{3.83}$$

Hence $\boldsymbol{\lambda}$ is achievable by a local policy $\tilde{\pi}$, so $\boldsymbol{\lambda} \in \Lambda^{\text{loc}}$. □

Corollary 1. *At $\text{INR} = \infty$, the stability region decomposes as an intersection of per-node regions:*

$$\Lambda^{\text{loc}}|_{\text{INR}=\infty} = \bigcap_{n \in \mathcal{B}} \Lambda_n^{\text{HD}},\tag{3.84}$$

where the per-node region is: for the root, $\Lambda_r^{\text{HD}} := \{\boldsymbol{\lambda} : [\lambda_\ell]_{\ell \in \mathcal{L}_r} \in \text{Conv}(\mathcal{C}_r^{\text{IDLE}})\}$, and for each

non-root n ,

$$\Lambda_n^{\text{HD}} := \left\{ \boldsymbol{\lambda} : \frac{[\lambda_\ell]_{\ell \in \mathcal{L}_n}}{1 - \lambda_{b_n}/\bar{R}_{b_n}(0)} \in \sum_{\boldsymbol{\mu}_n} \pi_{\boldsymbol{\mu}_n} \text{Conv}(\mathcal{C}_{\boldsymbol{\mu}_n}^{\text{IDLE}}) \right\}, \quad (3.85)$$

with $\mathcal{C}_{\boldsymbol{\mu}_n}^{\text{IDLE}} := \{(\mu_\ell \cdot s_\ell)_{\ell \in \mathcal{L}_n} : \mathbf{s} \in \mathcal{S}_n^{\text{IDLE}}\}$ and $\bar{R}_{b_n}(0) := \sum_{\boldsymbol{\mu}_n} \pi_{\boldsymbol{\mu}_n} R_{b_n}(0, \boldsymbol{\mu}_n)$. This coincides with the region $\Lambda_{\mathcal{P}} = \bigcap_n \Lambda_n$ of [8, eq. (20)].

Proof. At $\text{INR} = \infty$: $R_{b_n}(m, \boldsymbol{\mu}_n) = 0$ for all $m \geq 1$ and all $\boldsymbol{\mu}_n$.

Step 1: BH constraint forces $\alpha_n(0)$. From (LP-BH): $\sum_m \sum_{\boldsymbol{\mu}_n} \alpha_n(m, \boldsymbol{\mu}_n) R_{b_n}(m, \boldsymbol{\mu}_n) \geq \lambda_{b_n}$. Since $R_{b_n}(m, \boldsymbol{\mu}_n) = 0$ for $m \geq 1$, only the $m = 0$ term survives: $\sum_{\boldsymbol{\mu}_n} \alpha_n(0, \boldsymbol{\mu}_n) R_{b_n}(0, \boldsymbol{\mu}_n) \geq \lambda_{b_n}$, i.e., $\rho_n \cdot \bar{R}_{b_n}(0) \geq \lambda_{b_n}$, where $\bar{R}_{b_n}(0) := \sum_{\boldsymbol{\mu}_n} \pi_{\boldsymbol{\mu}_n} R_{b_n}(0, \boldsymbol{\mu}_n)$.

Step 2: No active links when $\theta_n = 0$. From (LP-TX) at $m = 0$: $\sum_\ell y_{n,\ell}^{\text{T},0} = 0 \cdot \alpha_n(0) = 0$, so $y_{n,\ell}^{\text{T},0} = 0$ for all ℓ .

Step 3: Modes $\theta_n = m \geq 1$ are suboptimal. When $\theta_n = m \geq 1$: panel P_n^1 is pinned to BH-DL RX yet $R_{b_n}(m, \boldsymbol{\mu}_n) = 0$, so the node loses one panel without BH benefit. This is strictly worse than mode ϕ (all P_n panels free). Hence any LP-feasible solution can be improved by setting $\alpha_n(m) = 0$ for $m \geq 1$, which implies $y_{n,\ell}^{\text{T},m} = 0$ for $m \geq 1$.

Step 4: All service in mode ϕ . The available mode- ϕ time is $\alpha_n^{\text{IDLE}} = 1 - \alpha_n(0)$. All remaining LP constraints involve only $y_{n,\ell}^{\text{IDLE}}$:

- (LP-AC): $\bar{\mu}_\ell y_{n,\ell}^{\text{IDLE}} \geq \lambda_\ell$ for each $\ell \in \mathcal{L}_n$,
- (LP-PNL^{IDLE}): $\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{IDLE}} \leq \alpha_n^{\text{IDLE}}$ for each panel k .

Step 5: Reduction to convex hull condition. Define the conditional activation $x_{n,\ell}^{\text{IDLE}} := y_{n,\ell}^{\text{IDLE}}/\alpha_n^{\text{IDLE}}$. By (LP-PNL^{IDLE}): $\sum_{\ell \in \mathcal{L}_{n,k}} x_{n,\ell}^{\text{IDLE}} \leq 1$ for each panel k , so $\mathbf{x}^{\text{IDLE}} \in \text{Conv}(\mathcal{S}_n^{\text{IDLE}})$. The rate vector $(\bar{\mu}_\ell x_{n,\ell}^{\text{IDLE}})_{\ell \in \mathcal{L}_n} \in \text{Conv}(\mathcal{C}_n^{\text{IDLE}})$, and (LP-AC) requires each component to satisfy $\bar{\mu}_\ell x_{n,\ell}^{\text{IDLE}} \geq \lambda_\ell/\alpha_n^{\text{IDLE}}$. The tightest condition is at the minimum BH time $\rho_n = \lambda_{b_n}/\bar{R}_{b_n}(0)$, yielding $\alpha_n^{\text{IDLE}} = 1 - \lambda_{b_n}/\bar{R}_{b_n}(0)$ and the requirement (3.85).

Step 6: Per-node independence. The parent-child coupling (LP-PC) requires $\rho_c = \lambda_{b_c}/\bar{R}_{b_c}(0)$, which is fully determined by $\boldsymbol{\lambda}$. Hence each node's condition (3.85) depends only

on λ and not on other nodes' LP variables. The overall LP is feasible if and only if (3.85) holds at every node simultaneously, giving the intersection (3.84). $\square$

Theorem 4. *If $\lambda \notin \Lambda^{\text{loc}}$, then no local policy stabilizes the system.*

Proof. We prove the contrapositive: if a stable local policy π exists, then $\lambda \in \Lambda^{\text{loc}}$.

Step 1: Time-averaged fractions. By Definition 17 and Lemma 6, the fractions $(\rho_n, \alpha_n(m, \mu_n), y_{n,\ell}^{\text{T},m}(\mu_n), y_{n,\ell}^{\text{IDLE}}(\mu_n))$ are well-defined and non-negative. We verify these satisfy every LP constraint.

Step 2: Independence constraint (LP-IND). Since π is local, parent's BH activation is independent of child's channel: $\mathbb{P}(\theta_n \neq \phi \mid \mu_n) = \mathbb{P}(\theta_n \neq \phi) = \rho_n$. Hence:

$$\sum_m \alpha_n(m, \mu_n) = \mathbb{P}(\theta_n \neq \phi, \mu_n) = \rho_n \pi_{\mu_n}. \quad (3.86)$$

Step 3: BH-DL service (LP-BH). The time-averaged BH-DL service rate is

$$\bar{S}_{b_n} = \sum_m \sum_{\mu_n} R_{b_n}(m, \mu_n) \alpha_n(m, \mu_n), \quad (3.87)$$

by decomposing over the joint events $\{\theta_n = m, \mu_n\}$. Mean stability requires $\bar{S}_{b_n} \geq \lambda_{b_n}$: if $\bar{S}_{b_n} < \lambda_{b_n}$, then $\frac{1}{T} \sum_t \mathbb{E}[\Delta q_{b_n}(t)] \geq \lambda_{b_n} - \bar{S}_{b_n} > 0$, so $\mathbb{E}[q_{b_n}(T)]/T \rightarrow \infty$, contradicting mean stability.

Step 4: Access-link service (LP-AC). For each $\ell \in \mathcal{L}_n^{\text{AC}}$:

$$\bar{S}_\ell = \lim_T \frac{1}{T} \sum_t \mathbb{E}[\mu_\ell(t) s_\ell(t)] = \sum_m \sum_{\mu_n} \mu_\ell y_{n,\ell}^{\text{T},m}(\mu_n) + \sum_{\mu_n} \mu_\ell y_{n,\ell}^{\text{IDLE}}(\mu_n), \quad (3.88)$$

by decomposing the expectation over the joint events $\{\theta_n=m, \mu_n(t)=\mu_n, s_\ell=1\}$, where μ_ℓ takes value μ_ℓ when $\mu_n(t) = \mu_n$. Mean stability requires $\bar{S}_\ell \geq \lambda_\ell$ by the same drift argument: if $\bar{S}_\ell < \lambda_\ell$, then $\mathbb{E}[q_\ell^f(T)]/T \rightarrow \infty$ for some flow f using link ℓ .

Step 5: Per-panel constraints (LP-PNL^T) and (LP-PNL^{IDLE}). The physical constraint

$\sum_{\ell \in \mathcal{L}_{n,k}} s_\ell(t) \leq 1$ holds in every slot. Restricting to the joint event $\{\theta_n = m, \boldsymbol{\mu}_n(t) = \boldsymbol{\mu}_n\}$:

$$\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) \leq \alpha_n(m, \boldsymbol{\mu}_n). \quad (3.89)$$

Similarly, $\sum_{\ell \in \mathcal{L}_{n,k}} y_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n) \leq (1 - \rho_n) \pi_{\boldsymbol{\mu}_n}$.

Step 6: TX count (LP-TX). When $\theta_n = m$, exactly m TX panels are active regardless of $\boldsymbol{\mu}_n$:

$$\sum_{\ell \in \mathcal{L}_n^{\text{TX}}} y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) = m \alpha_n(m, \boldsymbol{\mu}_n). \quad (3.90)$$

Step 7: Parent-child coupling (LP-PC). Child c enters Mode T iff parent p activates b_c .

Summing over parent's modes and channel states:

$$\sum_m \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{T},m}(\boldsymbol{\mu}_p) + \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{IDLE}}(\boldsymbol{\mu}_p) = \mathbb{P}(s_{b_c}=1) = \mathbb{P}(\theta_c \neq \phi) = \rho_c. \quad (3.91)$$

Step 8. Steps 2–7 show that the fractions satisfy every LP constraint. By Proposition 1, $\boldsymbol{\lambda} \in \Lambda^{\text{loc}}$. Taking the contrapositive: $\boldsymbol{\lambda} \notin \Lambda^{\text{loc}}$ implies no stable local policy exists. $\square$

Fluid Model. The stability proofs use the *fluid model* framework of Dai [6]. Consider a sequence of stochastic systems indexed by scaling parameter n , with initial condition $\mathbf{Q}^{(n)}(0)$ satisfying $\|\mathbf{Q}^{(n)}(0)\| = n$. Define the *fluid-scaled queue process*:

$$\bar{q}_\ell^{(n)}(\tau) := \frac{Q_\ell^{(n)}(\lfloor n\tau \rfloor)}{n}, \quad \tau \geq 0. \quad (3.92)$$

A *fluid model solution* $\bar{\mathbf{q}}(\tau)$ is any subsequential limit $\bar{q}_\ell(\tau) := \lim_{k \rightarrow \infty} \bar{q}_\ell^{(n_k)}(\tau)$ (pointwise in τ). Under mild regularity conditions (bounded second moments of arrivals and service), fluid model solutions exist and are Lipschitz continuous [6, Theorem 4.1]. When $\bar{q}_\ell(\tau) > 0$, the fluid queue evolves as:

$$\dot{\bar{q}}_\ell(\tau) = \bar{a}_\ell(\tau) - \bar{r}_\ell(\tau), \quad (3.93)$$

where $\bar{a}_\ell(\tau)$ is the fluid arrival rate and $\bar{r}_\ell(\tau)$ is the fluid service rate at time τ , both determined by the scheduling policy. The dot denotes the time derivative $d/d\tau$. When $\bar{q}_\ell(\tau) = 0$, the queue remains at zero as long as $\bar{a}_\ell \leq \bar{r}_\ell$.

Lemma 7. (Dai [6, Theorem 4.2].) *If every fluid model solution satisfies $\bar{\mathbf{q}}(T) = \mathbf{0}$ for some finite T (depending on the initial condition $\bar{\mathbf{q}}(0)$), then the stochastic network is positive recurrent.*

Theorem 5. *If $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$, then there exists a local policy that stabilizes the system.*

Proof. We construct an explicit stationary randomized local policy and show stability via the Dai fluid model. *Step 1: LP solution with slack.* Since $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$, there exists $\delta > 0$ such that $\boldsymbol{\lambda} + \delta \mathbf{1} \in \Lambda^{\text{loc}}$. By Proposition 1, the LP is feasible at $\boldsymbol{\lambda} + \delta \mathbf{1}$. Let $(\rho_n^*, \alpha_n^*(m, \boldsymbol{\mu}_n), y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n), y_{n,\ell}^{\star\phi}(\boldsymbol{\mu}_n))$ denote a feasible solution. It satisfies:

$$\sum_m \sum_{\boldsymbol{\mu}_n} \alpha_n^*(m, \boldsymbol{\mu}_n) R_{b_n}(m, \boldsymbol{\mu}_n) \geq \lambda_{b_n} + \delta, \quad (3.94)$$

$$\sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) + \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\star\phi}(\boldsymbol{\mu}_n) \geq \lambda_\ell + \delta, \quad \forall \ell \in \mathcal{L}_n^{\text{AC}}, \quad (3.95)$$

$$\sum_m \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\text{T},m}(\boldsymbol{\mu}_p) + \sum_{\boldsymbol{\mu}_p} y_{p,b_c}^{\star\phi}(\boldsymbol{\mu}_p) = \rho_c^*. \quad (3.96)$$

Step 2: Channel-conditional policy construction. For each n, m , and $\boldsymbol{\mu}_n$ with $\alpha_n^*(m, \boldsymbol{\mu}_n) > 0$, define

$$z_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n) := \frac{y_{n,\ell}^{\text{T},m}(\boldsymbol{\mu}_n)}{\alpha_n^*(m, \boldsymbol{\mu}_n)}. \quad (3.97)$$

By (LP-PNL^T) and (LP-TX), $\mathbf{z}_n^{\text{T},m}(\boldsymbol{\mu}_n) \in \text{Conv}(\mathcal{S}_n^{\text{T},m})$ for each $\boldsymbol{\mu}_n$. By Theorem 3, each decomposes into a distribution $\{p_j^{\text{T},m}(\boldsymbol{\mu}_n)\}$ over at most $|\mathcal{L}_n|+1$ feasible schedules. Similarly for $z_{n,\ell}^{\text{IDLE}}(\boldsymbol{\mu}_n)$. Construct policy $\tilde{\pi}$: at each slot t , independently of $Q(t)$:

- (a) *Root r* : observe $\boldsymbol{\mu}_r(t)$, draw $\mathbf{s}_r \sim \{p_j^{\text{IDLE}}(\boldsymbol{\mu}_r(t))\}$.
- (b) *Non-root n* : if parent activated b_n (Mode T), observe $\boldsymbol{\mu}_n(t)$, draw m with probability $\alpha_n^*(m, \boldsymbol{\mu}_n(t))/(\rho_n^* \pi_{\boldsymbol{\mu}_n(t)})$, then draw $\mathbf{s}_n \sim \{p_j^{\text{T},m}(\boldsymbol{\mu}_n(t))\}$. If Mode ϕ , draw

$$\mathbf{s}_n \sim \{p_j^{\text{IDLE}}(\boldsymbol{\mu}_n(t))\}.$$

This policy is local: each node uses only its mode (from parent) and its own channel $\boldsymbol{\mu}_n(t)$. By (LP-IND), the total Mode T fraction is ρ_n^* regardless of $\boldsymbol{\mu}_n$, preserving the independence property. Coupling (3.96) ensures consistency across the tree. *Step 3: Per-link service rates under $\tilde{\pi}$.* For BH-DL at non-root n :

$$\bar{S}_{b_n} = \sum_m \sum_{\boldsymbol{\mu}_n} \alpha_n^*(m, \boldsymbol{\mu}_n) R_{b_n}(m, \boldsymbol{\mu}_n) \geq \lambda_{b_n} + \delta. \quad (3.98)$$

For each access link $\ell \in \mathcal{L}_n^{\text{AC}}$:

$$\bar{S}_\ell = \sum_m \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\star T,m}(\boldsymbol{\mu}_n) + \sum_{\boldsymbol{\mu}_n} \mu_\ell y_{n,\ell}^{\star \phi}(\boldsymbol{\mu}_n) \geq \lambda_\ell + \delta. \quad (3.99)$$

Step 4: Stability via fluid model. Under $\tilde{\pi}$, the scheduling decisions are i.i.d. across slots and independent of $Q(t)$. The offered service process $\{\mu_\ell(t)s_\ell(t)\}_t$ is i.i.d. (since channels are i.i.d. and $\mathbf{s}(t)$ is drawn from a channel-conditional but Q -independent distribution). The fluid model replaces stochastic arrivals and service with their means. Each fluid queue satisfies $\dot{\bar{q}} = \lambda - \bar{S} \cdot \mathbf{1}\{\bar{q} > 0\}$ with $\bar{S} \geq \lambda + \delta$. *Hierarchical draining.* By induction on tree depth:

- (i) *Root BH-DL queues* ($\bar{q}_{b_c}$): $\dot{\bar{q}}_{b_c} = \lambda_{b_c} - \bar{S}_{b_c} \leq -\delta$ when $\bar{q}_{b_c} > 0$. Drain by time $T_0 \leq \max_c \bar{q}_{b_c}(0)/\delta$.
- (ii) *Post-drain*: $\bar{q}_{b_c} = 0$ forces outflow $= \lambda_{b_c}$. Each downstream link receives arrival rate λ_ℓ .
- (iii) *Depth d* : once depth- $(d-1)$ BH queues drain, each depth- d queue has arrival λ and service $\bar{S} \geq \lambda + \delta$. Drains in finite time.

All fluid queues reach zero in finite time. By Lemma 7, $\tilde{\pi}$ is stable. *Applicability of Lemma 7.* Our construction satisfies the regularity conditions of [6, Theorem 4.2]: (i) arrivals have bounded second moments; (ii) the service set is finite; (iii) $\{\mu_\ell(t)s_\ell(t)\}_t$ is i.i.d. across slots;

(iv) BH-DL forwarding uses fixed proportional splitting. *Step 5: Conclusion.* The stationary local policy $\tilde{\pi}$ stabilizes the system for any $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$. $\square$

Stability of FD-LocalMW-BP

We prove stability of Algorithm 3 by induction on tree depth and fluid model.

Estimation of backhaul rate Write $\bar{\mu}_{b_c} := \mathbb{E}[R_{b_c}(m_c^*, \boldsymbol{\mu}_c) \mid s_{b_c}^{\text{BP}} = 1]$ for the true conditional mean under back-pressure at load $\boldsymbol{\lambda}$.

Proposition 2. *Since $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}}) \subseteq \text{int}(\Lambda)$, back-pressure stabilizes $\boldsymbol{\lambda}$ in the simulator (Lemma 2), and the conditional mean*

$$\bar{\mu}_{b_c} := \mathbb{E}_{\pi}[R_{b_c}(m_c^*, \boldsymbol{\mu}_c) \mid s_{b_c}^{\text{BP}} = 1] \quad (3.100)$$

is well-defined under the invariant distribution π of the simulated BP process at load $\boldsymbol{\lambda}$. Define the training estimation error

$$\epsilon_{\text{train}} := \max_c |\hat{\mu}_{b_c} - \bar{\mu}_{b_c}|. \quad (3.101)$$

Then $\epsilon_{\text{train}} \rightarrow 0$ a.s. as $T_{\text{est}} \rightarrow \infty$.

Proof. Under BP at $\boldsymbol{\lambda} \in \text{int}(\Lambda)$, the joint process $(\mathbf{Q}(t), \boldsymbol{\mu}(t))$ is a positive recurrent Markov chain with invariant distribution π . Since BP's scheduling decision $\mathbf{s}^{\text{BP}}(t)$ is a deterministic function of $(\mathbf{Q}(t), \boldsymbol{\mu}(t))$, the quantities $R_{b_c}(m_c^*(t), \boldsymbol{\mu}_c(t))$ and $s_{b_c}^{\text{BP}}(t)$ are measurable functions of the chain's state.

Define the bounded functions

$$f_1(t) := R_{b_c}(m_c^*(t), \boldsymbol{\mu}_c(t)) \cdot \mathbf{1}\{s_{b_c}^{\text{BP}}(t) = 1\}, \quad (3.102)$$

$$f_2(t) := \mathbf{1}\{s_{b_c}^{\text{BP}}(t) = 1\}. \quad (3.103)$$

By the Birkhoff ergodic theorem for positive recurrent Markov chains, time averages converge

to π -expectations almost surely:

$$\frac{1}{T} \sum_{t=0}^{T-1} f_1(t) \rightarrow \mathbb{E}_\pi[f_1] = \mathbb{E}_\pi[R_{b_c} \cdot \mathbf{1}\{s_{b_c} = 1\}], \quad \frac{1}{T} \sum_{t=0}^{T-1} f_2(t) \rightarrow \mathbb{E}_\pi[f_2] = \mathbb{P}_\pi(s_{b_c} = 1). \quad (3.104)$$

Since BP stabilizes λ with $\lambda_{b_c} > 0$, the BH-DL link b_c is activated with positive frequency: $\mathbb{P}_\pi(s_{b_c} = 1) > 0$. The estimate $\hat{\mu}_{b_c}$ is the ratio of these two averages:

$$\hat{\mu}_{b_c} = \frac{\sum_t f_1(t)}{\sum_t f_2(t)} = \frac{(1/T) \sum_t f_1(t)}{(1/T) \sum_t f_2(t)} \rightarrow \frac{\mathbb{E}_\pi[R_{b_c} \cdot \mathbf{1}\{s_{b_c} = 1\}]}{\mathbb{P}_\pi(s_{b_c} = 1)} = \mathbb{E}_\pi[R_{b_c} \mid s_{b_c} = 1] = \bar{\mu}_{b_c} \quad \text{a.s.} \quad (3.105)$$

The convergence of the ratio follows from the almost sure convergence of numerator and denominator, with the denominator bounded away from zero. Since $|\mathcal{C}|$ is finite, $\epsilon_{\text{train}} = \max_c |\hat{\mu}_{b_c} - \bar{\mu}_{b_c}| \rightarrow 0$ a.s. $\square$

Definition 18. For a non-root node n with $\lambda_{b_n} > 0$, let

$$\bar{R}_q^{(n)} := \frac{\int_0^T R_{b_n}(m^*(\tau), \boldsymbol{\mu}_n(\tau)) \bar{q}_{b_n}(\tau) \bar{u}_{b_n}(\tau) d\tau}{\int_0^T \bar{q}_{b_n}(\tau) \bar{u}_{b_n}(\tau) d\tau} \quad (3.106)$$

be the queue-weighted average of the backhaul rate, where $m^*(\tau)$ is the transmit count the child selects in its Mode- T program, and define the *estimation gap*

$$g_n := \epsilon_{\text{train}} + (\hat{\mu}_{b_n} - \bar{R}_q^{(n)})^+. \quad (3.107)$$

FD-LocalMW-BP never sees the rate backhaul link delivers, because that rate depends on how many panels the child transmits on, a local choice the parent cannot observe, so the parent acts on the $\hat{\mu}_{b_n}$ instead. This quantity is defined to charge the throughput this substitution costs. When the parent under-uses a link that could carry more, which does not threaten stability, $\hat{\mu}_{b_n}$ is smaller than the $\bar{R}_q^{(n)}$. But, when the parent over estimates backhaul link as if the link ran at $\hat{\mu}_{b_n}$ while delivers only $\bar{R}_q^{(n)}$, then the stability region can be reduced. Therefore, for arrival rates falling within the estimation gap of the local stability region, such

over-allocation does not occur, thereby guaranteeing stability.

Definition 19. Let $M_n(\tau) := \sum_{\ell \in \mathcal{L}_n} \lambda_\ell \bar{q}_\ell(\tau) + \sum_c \lambda_{b_c} \bar{q}_{b_c}(\tau)$ be the $\boldsymbol{\lambda}$ -weighted backlog at node n , and let σ be the policy π^λ of Step 1 restricted to the mode and channel $(\theta_n(\tau), \boldsymbol{\mu}_n(\tau))$ realized by FD-LocalMW-BP, with time-averaged service $\bar{S}_\ell^\sigma$ on access link ℓ and activation $u_{b_c}^\sigma$ on child-backhaul link b_c . The *operating gap* is

$$\zeta_n := \left(1 - \liminf_{T \rightarrow \infty} \frac{\int_0^T (\sum_\ell \bar{q}_\ell \bar{S}_\ell^\sigma + \sum_c \bar{q}_{b_c} \bar{\mu}_{b_c} u_{b_c}^\sigma) d\tau}{\int_0^T M_n d\tau} \right)^+, \quad (3.108)$$

Even with

Definition 20. For a reference load $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$, the *relative gap* is

$$\epsilon^*(\boldsymbol{\lambda}) := \sup \max_n \left(\zeta_n + \max_{c: \lambda_{b_c} > 0} \frac{g_c}{\lambda_{b_c}} \right), \quad (3.109)$$

the supremum taken over all FD-LocalMW-BP fluid trajectories at loads $(1-\epsilon)\boldsymbol{\lambda}$. When a node lights up a child link, the added self-interference lowers the rate of the parent backhaul it receives on, so serving children and receiving from the parent trade off against each other even when the estimate is exact. The operating gap is defined to charge the throughput a node loses when local policy resolves that tradeoff from its own queues rather than under global coordination.

Theorem 6. Let $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$ be the reference load used for BP estimation in the simulator (Phase 1). Then FD-LocalMW-BP mean-stabilizes $\tilde{\boldsymbol{\lambda}} = (1-\epsilon)\boldsymbol{\lambda}$ for every $\epsilon > \epsilon^*(\boldsymbol{\lambda})$.

Proof. Write $\tilde{\boldsymbol{\lambda}} = (1-\epsilon)\boldsymbol{\lambda}$ with $\epsilon > \epsilon^*(\boldsymbol{\lambda})$.

Step 1. Phase 1 runs BP at the reference load $\boldsymbol{\lambda} \in \text{int}(\Lambda^{\text{loc}})$, which BP stabilizes by throughput-optimality (Lemma 2). The long-run frequencies with which this run plays each schedule define a stationary randomized policy π^λ , and stability makes its mean service cover the load, $\bar{S}_\ell^\lambda \geq \lambda_\ell$ for all ℓ and $\bar{\mu}_{b_n} u_{b_n}^\lambda \geq \lambda_{b_n}$ for all $n \neq r$. The estimate $\hat{\mu}_{b_n}$ is the Phase 1

sample average of the conditional backhaul rate $\bar{\mu}_{b_n}$ of π^λ , so $|\bar{\mu}_{b_n} - \hat{\mu}_{b_n}| \leq \epsilon_{\text{train}}$ by the law of large numbers.

Step 2. In mode $\theta_n(\tau)$ and channel $\boldsymbol{\mu}_n(\tau)$, FD-LocalMW-BP solves the max-weight problem over $S_n(\theta_n, \boldsymbol{\mu}_n)$. In Mode T its objective is

$$\underbrace{R_{b_n}(m) \bar{q}_{b_n}}_{\text{parent-BH}} + \underbrace{\sum_{\ell \in \mathcal{L}_n} \mu_\ell \bar{q}_\ell s_\ell}_{\text{child access}} + \underbrace{\sum_c \hat{\mu}_{b_c} \bar{q}_{b_c} s_{b_c}}_{\text{child-BH}}; \quad (3.110)$$

In Mode ϕ the parent-BH term is absent.

The reference policy is $\sigma(\tau) := \pi^\lambda(\cdot \mid \theta_n(\tau), \boldsymbol{\mu}_n(\tau)) \in \text{Conv}(S_n(\theta_n(\tau), \boldsymbol{\mu}_n(\tau)))$, the Step 1 policy in the same mode and channel. Since FD-LocalMW-BP maximizes the local weight sum over $S_n(\theta_n, \boldsymbol{\mu}_n)$, its value is at least that of σ ; in Mode T ,

$$R_{b_n}(m^*) \bar{q}_{b_n} + \sum_\ell \bar{q}_\ell \bar{r}_\ell + \sum_c \hat{\mu}_{b_c} \bar{q}_{b_c} \bar{u}_{b_c} \geq R_{b_n}(m^\sigma) \bar{q}_{b_n} + \sum_\ell \bar{q}_\ell \bar{S}_\ell^\sigma + \sum_c \hat{\mu}_{b_c} \bar{q}_{b_c} u_{b_c}^\sigma. \quad (3.111)$$

When $\bar{q}_{b_n}$ is zero, (3.111) reduces to

$$\sum_{\ell \in \mathcal{L}_n} \bar{q}_\ell \bar{r}_\ell + \sum_c \hat{\mu}_{b_c} \bar{q}_{b_c} \bar{u}_{b_c} \geq \sum_{\ell \in \mathcal{L}_n} \bar{q}_\ell \bar{S}_\ell^\sigma + \sum_c \hat{\mu}_{b_c} \bar{q}_{b_c} u_{b_c}^\sigma. \quad (3.112)$$

Step 3: Induction on depth. We induct on depth d , showing that some finite T_d makes every fluid queue at depth $\leq d$ vanish for $\tau \geq T_d$. As the tree has finite depth D , all fluid queues then vanish by $T_D < \infty$, and by Lemma 7, the network is stable. Both cases run the same drift for a node n from a time τ_0 at which (3.112) holds and the arrivals obey $\bar{a}_\ell \leq \tilde{\lambda}_\ell$ and $\bar{a}_{b_c} \leq \tilde{\lambda}_{b_c}$.

Base case ($d = 0$). The root has no parent, so (3.112) holds for all τ , and its external arrivals satisfy $\bar{a}_\ell \leq \tilde{\lambda}_\ell$ and $\bar{a}_{b_c} \leq \tilde{\lambda}_{b_c}$. Take $\tau_0 = 0$.

Inductive step. Assume every fluid queue at depth $\leq d$ is zero for $\tau \geq T_d$, and let n have depth $d+1$ with parent p . The queue $\bar{q}_{b_n}$ resides at p , which has depth $\leq d$. By the

hypothesis, $\bar{q}_{b_n}(\tau) = 0$ for $\tau \geq T_d$ and (3.111) reduces to (3.112). Since the length of $\bar{q}_{b_n}(\tau)$ does not change, arrival rate should be less or equal to service rate; otherwise, $\bar{q}_{b_n}(\tau) \neq 0$.

$$\bar{a}_{b_n}(\tau) - \bar{R}_{b_n}(\tau) \bar{u}_{b_n}(\tau) \leq 0, \quad \tau \geq T_d, \quad (3.113)$$

Because the parent keeps a separate backhaul queue $\bar{q}_{b_n}^{(\ell)}$ for each downstream access link $\ell \in \text{subtree}(n)$, (3.113) holds for each flow, and by flow conservation each access link inherits $\bar{a}_\ell(\tau) \leq \tilde{\lambda}_\ell$ and each child link $\bar{a}_{b_c}(\tau) \leq \tilde{\lambda}_{b_c}$ for $\tau \geq T_d$. Take $\tau_0 = T_d$.

Drift. Let $\bar{V}_n(\tau) := \sum_{\ell \in \mathcal{L}_n} \bar{q}_\ell^2 + \sum_c \bar{q}_{b_c}^2$. Where all fluid queues at n are positive,

$$\frac{1}{2} \frac{d\bar{V}_n}{d\tau} = \sum_\ell \bar{q}_\ell (\bar{a}_\ell - \bar{r}_\ell) + \sum_c \bar{q}_{b_c} (\bar{a}_{b_c} - \bar{R}_{b_c} \bar{u}_{b_c}). \quad (3.114)$$

For $\tau \geq \tau_0$, writing $\bar{R}_{b_c} \bar{u}_{b_c} = \hat{\mu}_{b_c} \bar{u}_{b_c} + (\bar{R}_{b_c} - \hat{\mu}_{b_c}) \bar{u}_{b_c}$, applying (3.112), and using $\hat{\mu}_{b_c} u_{b_c}^\sigma \geq \bar{\mu}_{b_c} u_{b_c}^\sigma - \epsilon_{\text{train}}$,

$$\begin{aligned} \frac{1}{2} \frac{d\bar{V}_n}{d\tau} \leq & \underbrace{\sum_\ell \bar{q}_\ell \bar{a}_\ell + \sum_c \bar{q}_{b_c} \bar{a}_{b_c}}_{(a)} - \underbrace{\left(\sum_\ell \bar{q}_\ell \bar{S}_\ell^\sigma + \sum_c \bar{q}_{b_c} \bar{\mu}_{b_c} u_{b_c}^\sigma \right)}_{(b)} + \underbrace{\epsilon_{\text{train}} \sum_c \bar{q}_{b_c} - \sum_c \bar{q}_{b_c} (\bar{R}_{b_c} - \hat{\mu}_{b_c}) \bar{u}_{b_c}}_{(c)}. \end{aligned} \quad (3.115)$$

Integrate over $[\tau_0, T]$ and bound the three groups.

(a): With $\bar{a}_\ell \leq \tilde{\lambda}_\ell = (1-\epsilon)\lambda_\ell$ and $\bar{a}_{b_c} \leq (1-\epsilon)\lambda_{b_c}$,

$$\int_{\tau_0}^T \left(\sum_\ell \bar{q}_\ell \bar{a}_\ell + \sum_c \bar{q}_{b_c} \bar{a}_{b_c} \right) d\tau \leq (1-\epsilon) \int_{\tau_0}^T M_n d\tau. \quad (3.116)$$

(b): By the definition (3.108) of ζ_n , for any $\delta > 0$ and for all T ,

$$\int_{\tau_0}^T \left(\sum_\ell \bar{q}_\ell \bar{S}_\ell^\sigma + \sum_c \bar{q}_{b_c} \bar{\mu}_{b_c} u_{b_c}^\sigma \right) d\tau \geq (1 - \zeta_n - \delta) \int_{\tau_0}^T M_n d\tau, \quad (3.117)$$

(c): The backhaul rate enters directly through $\bar{R}_q^{(c)}$ of (3.106) with $\bar{u}_{b_c} \leq 1$,

$$\epsilon_{\text{train}} \int_{\tau_0}^T \sum_c \bar{q}_{b_c} d\tau - \int_{\tau_0}^T \sum_c \bar{q}_{b_c} (\bar{R}_{b_c} - \hat{\mu}_{b_c}) \bar{u}_{b_c} d\tau \leq \sum_c g_c \int_{\tau_0}^T \bar{q}_{b_c} d\tau \leq \max_c \frac{g_c}{\lambda_{b_c}} \int_{\tau_0}^T M_n d\tau, \quad (3.118)$$

where the last step uses $\sum_c \lambda_{b_c} \bar{q}_{b_c} \leq M_n$.

Combining (3.116)–(3.118),

$$\bar{V}_n(T) - \bar{V}_n(\tau_0) \leq -2 \left(\epsilon - \zeta_n - \delta - \max_c \frac{g_c}{\lambda_{b_c}} \right) \int_{\tau_0}^T M_n d\tau = -2\kappa_n \int_{\tau_0}^T M_n d\tau, \quad (3.119)$$

with $\kappa_n := \epsilon - \zeta_n - \delta - \max_c (g_c/\lambda_{b_c})$. Since $\epsilon > \epsilon^*(\boldsymbol{\lambda})$ strictly, $\kappa_n > 0$.

Let $\lambda_{\min} := \min\{\lambda_i : \lambda_i > 0, i \text{ incident to } n\}$. Then $M_n \geq \lambda_{\min}(\sum_\ell \bar{q}_\ell + \sum_c \bar{q}_{b_c}) \geq \lambda_{\min} \sqrt{\bar{V}_n}$, so (3.119) yields $\frac{d}{d\tau} \sqrt{\bar{V}_n} \leq -\kappa_n \lambda_{\min}$ on $[\tau_0, T]$ wherever $\bar{V}_n > 0$, hence

$$\sqrt{\bar{V}_n(\tau)} \leq \left(\sqrt{\bar{V}_n(\tau_0)} - \kappa_n \lambda_{\min} (\tau - \tau_0) \right)^+ \quad (3.120)$$

and every fluid queue at n reaches zero by $\tau_0 + \sqrt{\bar{V}_n(\tau_0)}/(\kappa_n \lambda_{\min})$ and stays there. In the base case ($\tau_0 = 0$) this gives a finite T_0 ; in the inductive step ($\tau_0 = T_d$),

$$T_{d+1} \leq T_d + \frac{\sqrt{\bar{V}_n(T_d)}}{\kappa_n \lambda_{\min}} < \infty, \quad (3.121)$$

completing the induction. □

Remark 1. Since $\epsilon^*(\boldsymbol{\lambda})$ is usually positive, FD-LocalMW-BP cannot stabilize every rate in Λ^{loc} . To confirm that this gap is nonetheless small, we turn to simulation. The gaps are abstract, defined through the reference policy π^λ and a supremum over fluid trajectories, so we instead observe similar observable quantities. The integral defining $\bar{R}_q^{(n)}$ in (3.106) runs over the full horizon, which the simulation truncates to its finite run length, and the resulting finite-time average returns g_n through (3.107). Since it is hard to compute (3.108) through numerical analysis, it is estimated from the BP reference load $\boldsymbol{\lambda}$ that defines π^λ in Step 1.

Let $\hat{S}_\ell(\theta, \boldsymbol{\mu})$ and $\hat{u}_{b_c}(\theta, \boldsymbol{\mu})$ be the empirical service and activation of that run conditioned on mode and channel. Then, ζ_n can be approximately obtained as

$$\zeta_n \approx \left(1 - \frac{\int_0^T (\sum_\ell \bar{q}_\ell \hat{S}_\ell(\theta_n, \boldsymbol{\mu}_n) + \sum_c \bar{q}_{b_c} \hat{\mu}_{b_c} \hat{u}_{b_c}(\theta_n, \boldsymbol{\mu}_n)) d\tau}{\int_0^T M_n d\tau} \right)^+, \quad (3.122)$$

Chapter 4

Numerical Results

4.1 Simulation Setup

We evaluate the proposed scheduling algorithms on the FD-IAB tree shown in Fig. 4.1. The network consists of one IAB donor and ten non-donor IAB nodes arranged in a multi-hop tree. Each IAB node is equipped with $P = 4$ directional panels. Each panel can be assigned to at most one incident link in a slot, while different panels at the same IAB node may be used simultaneously for transmission and reception under full-duplex operation. Both access and backhaul links are modeled in downlink and uplink directions.

The backhaul and access link distances are randomly generated. For each backhaul link, the distance between the parent IAB and the child IAB is drawn uniformly from $[340, 440]$ m. For each access link, the UE distance from its serving IAB is drawn uniformly from $[10, 200]$ m. Each panel independently serves a uniformly random number of UEs from $[0, 5]$.

Packets arrive exogenously at the source queues according to i.i.d Poisson random process. The aggregated mean of uplink and downlink is denoted by λ Mbps. Unless otherwise stated, the uplink traffic ratio is set to 0.1, and the remaining traffic is downlink traffic. The mean is chosen to be the same for each UE.

The physical layer parameters are summarized in Table 4.1. We use carrier frequency

Table 4.1: Simulation parameters.

Parameter	Value
Carrier frequency	$f_c = 23$ GHz
Bandwidth	$B = 1$ GHz
Slot duration	$T_s = 125$ μ s
Packet size	$L_{\text{pkt}} = 100$ kbits
Noise power spectral density	-174 dBm/Hz
IAB receiver noise figure	5 dB
UE receiver noise figure	7 dB
Path-loss model	$32.4 + 21 \log_{10}(d) + 20 \log_{10}(f_c)$ dB
BH transmit power	24 dBm
BH beamforming gain	40 dB
AC-DL transmit power	24 dBm
AC-DL beamforming gain	30 dB
AC-UL transmit power	23 dBm
AC-UL beamforming gain	23 dB
Backhaul distance	Uniform in $[340, 440]$ m
Access distance	Uniform in $[10, 200]$ m

$f_c = 23$ GHz, bandwidth $B = 1$ GHz, slot duration $T_s = 125$ μ s, and packet size $L_{\text{pkt}} = 100$ kbits. The path loss at distance d meters is $\text{PL}_{\text{dB}}(d) = 32.4 + 21 \log_{10}(d) + 20 \log_{10}(f_c)$. The noise power spectral density is -174 dBm/Hz. The receiver noise figures are 5 dB for IAB receivers and 7 dB for UE receivers. For backhaul links, the transmit power is 24 dBm and the beamforming gain is 40 dB. For access downlink links, the transmit power is 24 dBm and the beamforming gain is 30 dB. For access uplink links, the UE transmit power is 23 dBm and the total beamforming gain is 23 dB. Following [18], we consider Rician fading for access links with K factor as 13 dB for LOS links and 6 dB for NLOS links. The fadings are independently generated in each slot. During the estimation phase, each flow source adds virtual packets to its real arrivals so that the effective arrival rate is 1.05 times the real rate. These virtual packets are routed and scheduled identically to real packets, so back-pressure activates the backhaul links slightly more often, but they are discarded at the destination and excluded from all reported metrics.

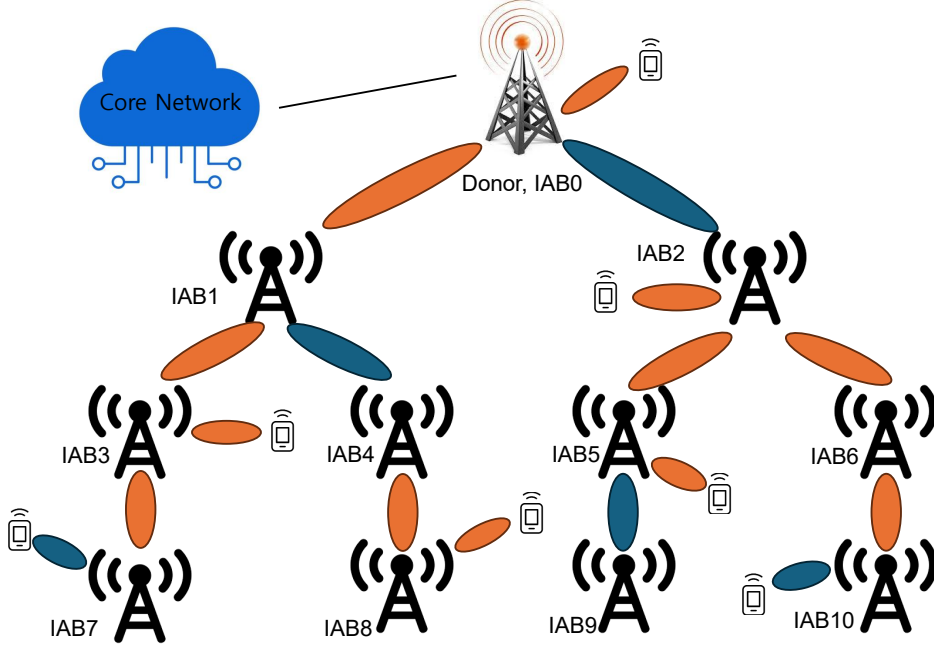


Figure 4.1: Simulated IAB topology.

Scheme	Duplex	Information	BH-DL rate used
MaxWeight	Full	Global	True rate $R_{b_n}(m)$
BP-FD	Full	Global	True rate $R_{b_n}(m)$
BP-HD [8]	Half	Global	No-SI rate $R_{b_n}(0)$
FD-LocalMW-BP	Full	Local	Estimated weight $\hat{\mu}_{b_n}$
Optimistic	Full	Local	No-SI rate $R_{b_n}(0)$
Pessimistic	Full	Local	Worst-case rate $R_{b_n}(P-1)$

Table 4.2: Schedulers compared in the evaluation.

Scheduling Policies for Comparison

We consider the following six policies for comparison.

Figure 4.2 plots the mean UE delay against the arrival rate λ , at three self-interference levels $\text{INR} \in \{-10, 0, +10\}$ dB. The half-duplex scheme BP-HD behaves the same at every INR, since it avoids self-interference, but becomes unstable beyond 130 Mbps, whereas the full-duplex schemes stay stable far past that point, an advantage that narrows but persists as INR rises. The three baselines and the proposed scheme reach stable regions comparable to those of the global schedulers, and they differ only in how often each activates the backhaul.

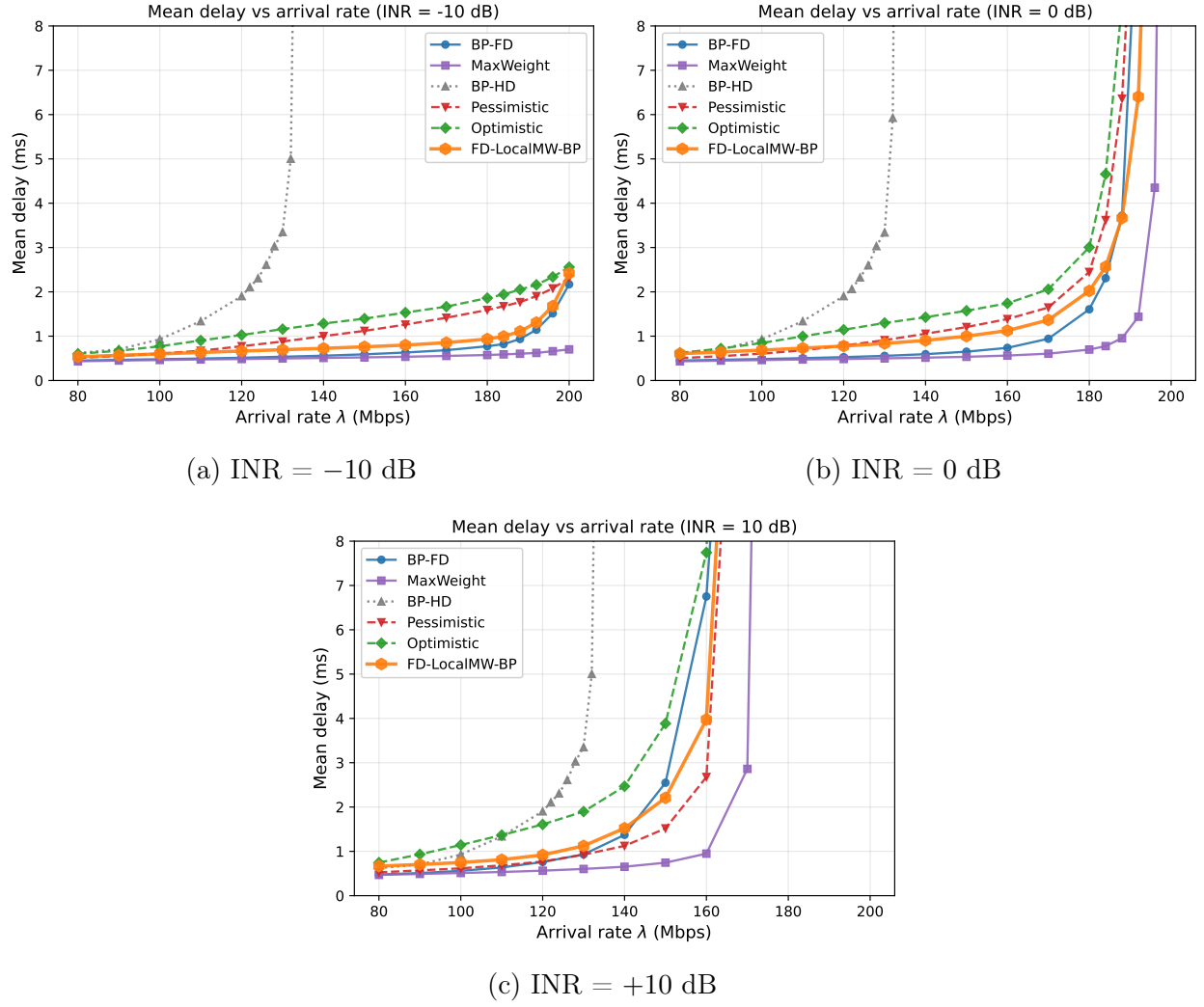


Figure 4.2: Mean delay versus arrival rate λ , at three self-interference levels.

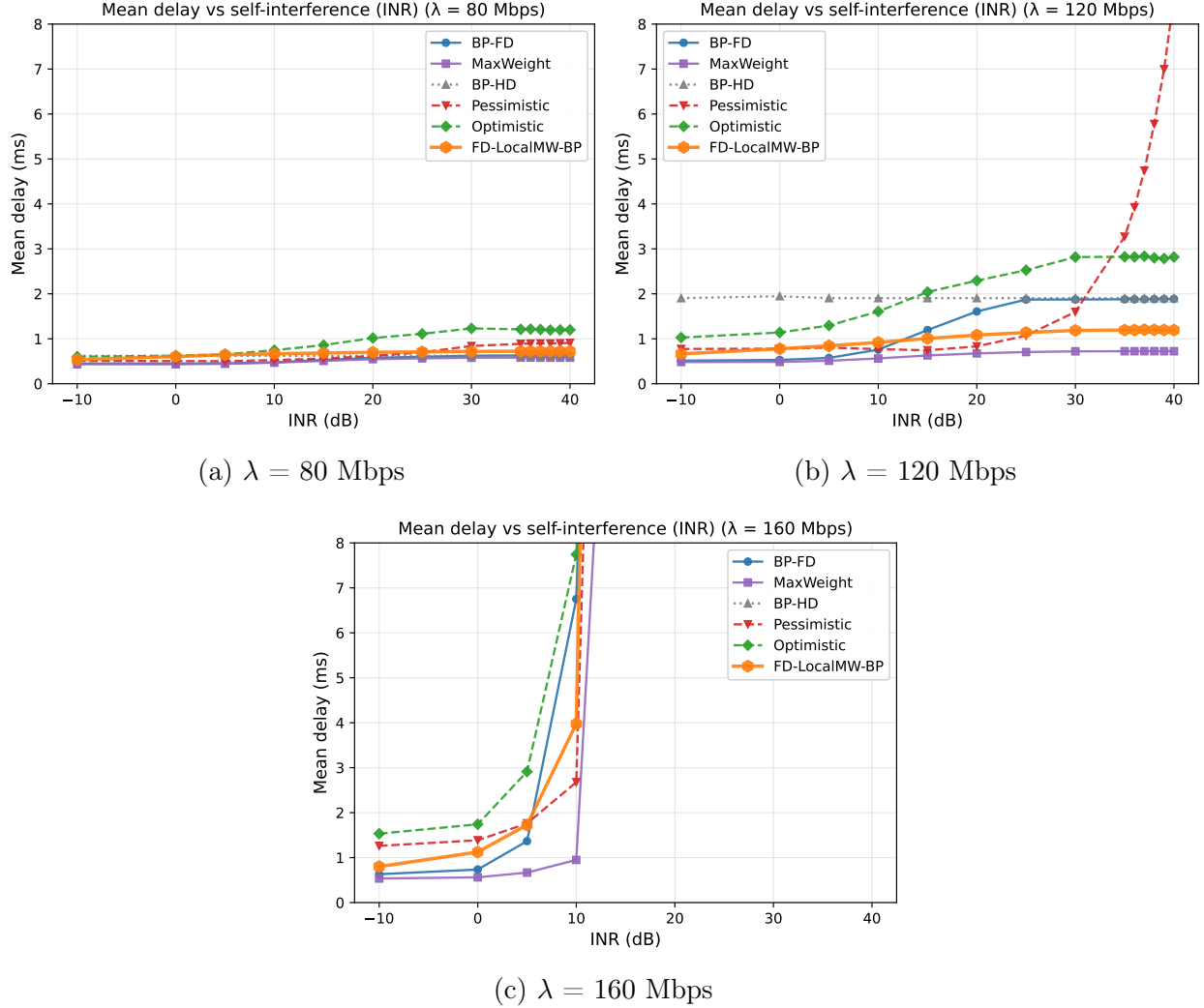
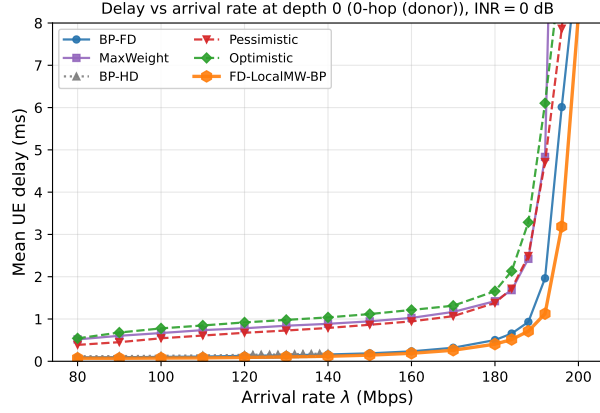


Figure 4.3: Mean delay versus self-interference (INR), at three arrival rates.

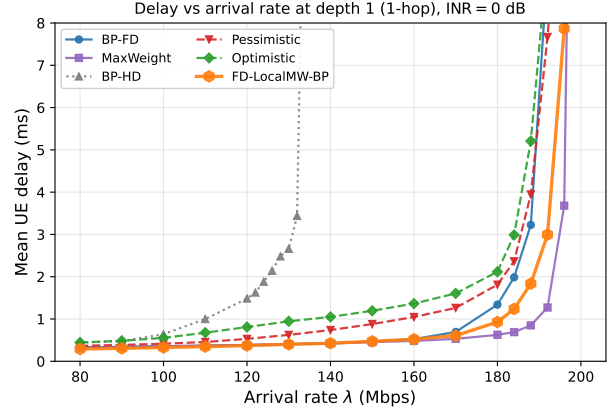
Optimistic assumes the highest rate and over-activates the backhaul, on a rate that the child's concurrent transmissions then drive below what the weight predicted and achieve low throughput and high delays. FD-LocalMW-BP instead activates the backhaul at about the frequency the child can absorb. As INR rises, the edge of every full-duplex scheme moves to a lower arrival rate, because each activation then drains little while still pinning a child panel under self-interference. At $\text{INR} = +10$ dB Pessimistic reaches a lower delay than FD-LocalMW-BP, since activation is rarely worth its cost at this low rate. This advantage is narrow and specific to this regime, and the next subsection examines the mechanism behind the inversion.

Figure 4.3 demonstrates the mean UE delay against the self-interference level INR for all seven schedulers, at three fixed arrival rates $\lambda \in \{80, 120, 160\}$ Mbps. At 80 Mbps, as the arrival rate is sufficiently small, even with high SI, network keeps stability and small delay. At 120 Mbps, as INR grows, BP-FD increasingly avoids self-interference by separating transmission and reception into different slots, so it behaves like a half-duplex node and its performance converges to that of BP-HD. Among the local schemes FD-LocalMW-BP is the most robust to self-interference, with its delay staying low and nearly flat and reaching only about 1.2 ms at 40 dB. On the other hand, delay of Pessimistic diverges since it always evaluates the backhaul-DL link at the worst-case self-interference level. As INR grows, the rate $R_{b_c}(P-1)$ it assigns to the link falls toward zero, so its backhaul weight shrinks and the parent almost stops activating the link. The packets waiting for the whole downstream subtree can then no longer be forwarded, the backhaul-DL queues keep growing, and the delay diverges. At 160 Mbps, as HD can not sustain such a large arrival rate, the network becomes unstable beyond 10 dB.

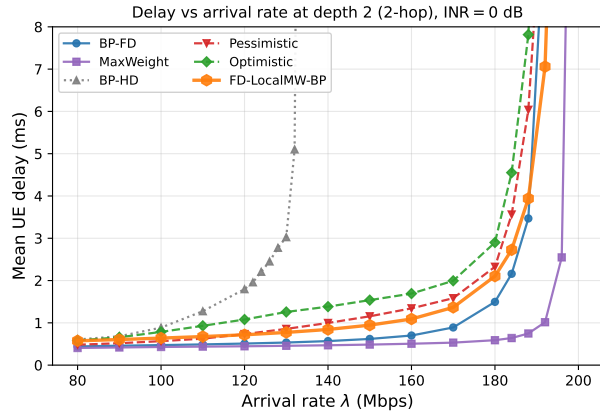
Figure 4.4 plots the mean UE delay against the arrival rate λ at $\text{INR} = 0$ dB for each tree depth 0–3, the number of backhaul hops from the donor. The minimum delay grows with depth since an end-to-end delay adds one queueing delay per hop and deeper UEs accumulate more hops. The half-duplex baseline BP-HD stays as low as the full-duplex schemes at the donor, which has no parent and serves its own UEs in every slot, but at depth 1 an inflexibility of HD cannot sustain even 130 Mbps. Among the full-duplex schemes only Pessimistic and FD-LocalMW-BP swap positions with depth, because activating the backhaul costs the child a serving slot and pays off only when its backhaul queue is large enough to be worth draining. At shallow depth the backhaul queues hold the aggregated traffic of large subtrees, so draining them helps the whole subtree and Pessimistic, which activates least, sits above FD-LocalMW-BP. At depth 3 the queue is small, so activating mostly pins the child to reception and delays its own UEs, and Pessimistic falls slightly below FD-LocalMW-BP. Optimistic activates most at every depth and stays highest throughout,



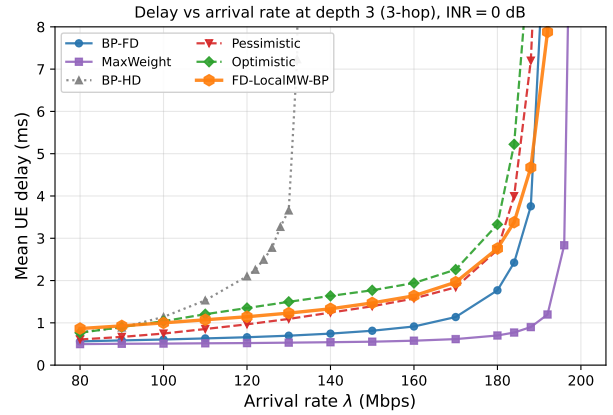
(a) Donor, 0-hop



(b) 1-hop



(c) 2-hop



(d) 3-hop

Figure 4.4: Mean UE delay versus arrival rate λ at $\text{INR} = 0$ dB, shown separately for each tree depth.

while FD-LocalMW-BP activates often enough to clear the large shallow queues yet not so often as to pin the lightly loaded depth-3 nodes.

Figure 4.5 shows the CDF of per-packet delay at $\lambda = 120$ Mbps and $\text{INR} = 0$ dB. The global schedulers MaxWeight and BP-FD have almost no tail, and FD-LocalMW-BP closely follows them, so its slowest packets are served nearly as quickly. The other local baselines instead develop a long tail, and the half-duplex baseline BP-HD has the heaviest tail by a wide margin since it cannot exploit FD operation. A heavy tail pulls the mean delay up, which is why the baselines report higher mean delays. This matches the mechanisms seen earlier, since Optimistic over-activates the backhaul on a rate the child cannot deliver

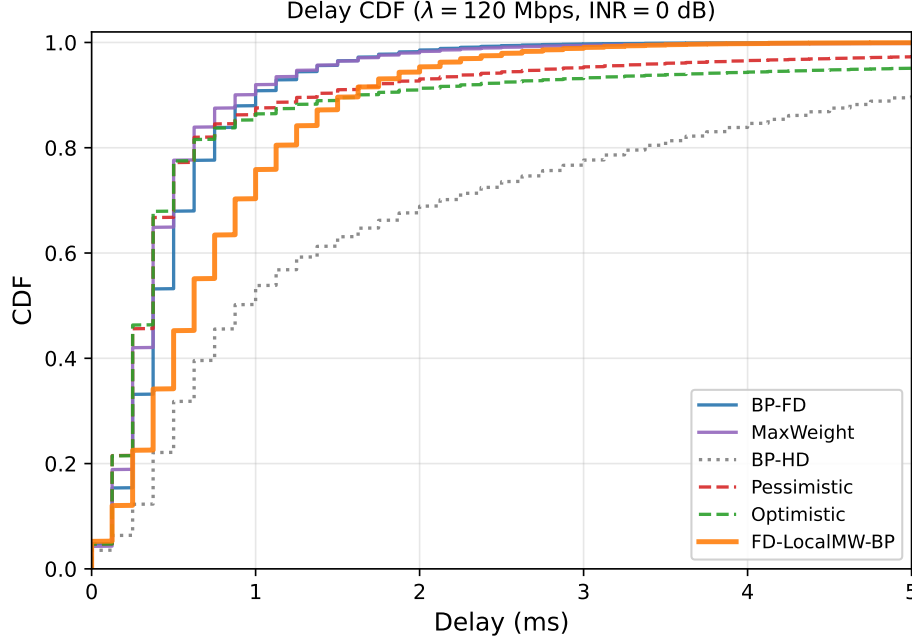


Figure 4.5: CDF of individual delay at $\lambda = 120$ Mbps and $\text{INR} = 0$ dB.

and Pessimistic under-activates it, so in both a fraction of packets wait through repeated misallocations and form the tail, while FD-LocalMW-BP keeps that fraction small using only local information.

Figure 4.6 shows, for each active backhaul-DL link with 120–200 Mbps, the difference between the estimated weight $\hat{\mu}_{b_n}$ and the realized backhaul-DL rate $\bar{R}_q^{(n)}$, where $\bar{R}_q^{(n)}$ is the rate the link actually delivers averaged over the active backhaul slots and weighted by the backhaul queue Q_{b_n} of each slot. The difference is negative for essentially every sample, so $\bar{R}_q^{(n)}$ exceeds the estimated weight almost everywhere. Only its positive part adds to the margin $\epsilon^*(\boldsymbol{\nu})$ of Theorem 6, so this contribution is negligible. The weighting by Q_{b_n} is the reason. The slots with the largest Q_{b_n} dominate the average, and on those slots the child holds its transmit count low, since the backhaul term $R_{b_n}(m)Q_{b_n}$ then dominates its panel objective and a lower transmit count protects the backhaul-DL reception from self-interference. A lower transmit count means a higher realized rate, so the large- Q_{b_n} slots that dominate $\bar{R}_q^{(n)}$ are exactly the high-rate slots, which pulls $\bar{R}_q^{(n)}$ above the estimated weight $\hat{\mu}_{b_n}$ measured under back-pressure.

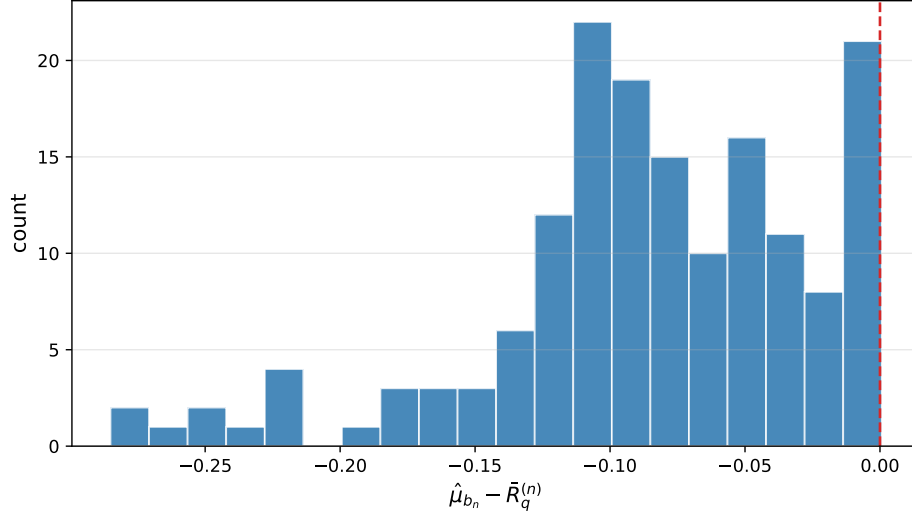


Figure 4.6: Distribution of the difference between the estimated weight $\hat{\mu}_{b_n}$ and the backhaul-DL rate $\bar{R}_q^{(n)}$ over all active backhaul-DL.

Figure 4.7 plots $\max_n \zeta_n$, estimated from the simulation through (3.122), against the arrival rate at three self-interference levels. Inside the local stability region the gap stays close to zero at every node and every self-interference level. The gap departs from zero only as the arrival rate approaches the boundary of that region, where the backhaul queue is so large that the backhaul is already activated almost every slot, leaving no further activation to draw on.

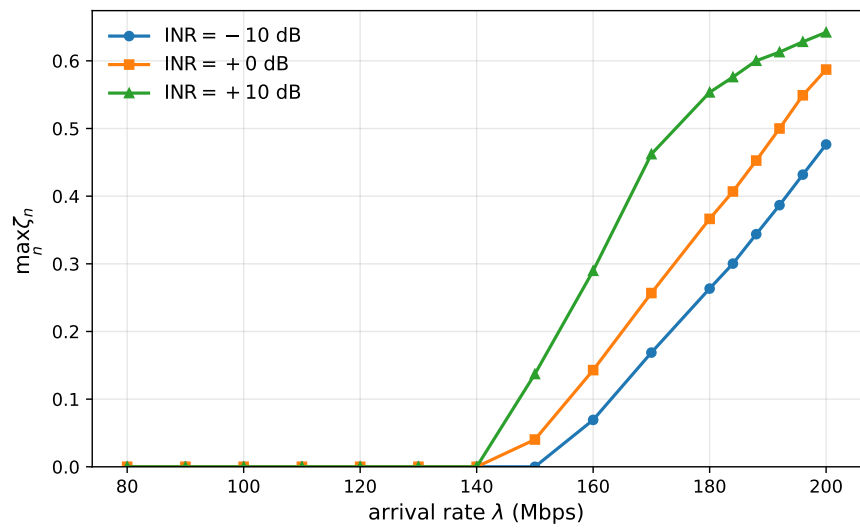


Figure 4.7: Approximated $\max_n \zeta_n$ of 3.122.

Chapter 5

Conclusion

In conclusion, this work first implemented backpressure scheduling and max-weight scheduling through a dynamic programming approach for the full-duplex IAB scheduling. Furthermore, we constructed a local schedule through learning with a global algorithm and analysed the stability region of the scheduling. We also empirically show that a local algorithm with low control overhead can achieve a comparable stability region. Since the scheduling-dependent weight is not restricted to a certain form in this work, we believe that it can be easily expanded to other forms.

Several possible future improvements are relevant. Firstly, the network structure can be generalised. Since this work used a tree structure, all routes of packets are uniquely determined. This can be expanded to a general directed acyclic graph (DAG) where multiple routes of packets are possible. In that problem, both routing and activation problems should simultaneously be considered. Another future research direction is constructing age-of-information (AoI)-based policy. Since the backpressure policy suffers from high delays due to a single flow limitation of a link, an AGI-based approach can enhance the delay much better.

Bibliography

- [1] 3GPP. NR; NR and NG-RAN overall description; stage-2. Technical Specification TS 38.300, 3rd Generation Partnership Project (3GPP), 2022. Section 4.7.1: Integrated Access and Backhaul.
- [2] Mamta Agiwal, Abhishek Roy, and Navrati Saxena. Next generation 5g wireless networks: A comprehensive survey. *IEEE Commun. surv. & tuts.*, 18(3):1617–1655, Feb. 2016.
- [3] Can Chen and Jinsong Gui. Hierarchical deep reinforcement learning-based resource allocation in dynamic time division duplex self-backhauling cellular networks. *IEEE Trans. Mob. Comput.*, pages 1–18, 2026. Early Access.
- [4] Gen-Huey Chen, MT Kuo, and JP Sheu. An optimal time algorithm for finding a maximum weight independent set in a tree. *BIT*, 28(2):353–356, 1988.
- [5] Tingjun Chen, Jelena Diakonikolas, Javad Ghaderi, and Gil Zussman. Hybrid scheduling in heterogeneous half-and full-duplex wireless networks. *IEEE/ACM Trans. Netw.*, 28(2):764–777, Apr. 2020.
- [6] Jim G Dai. On positive harris recurrence of multiclass queueing networks: a unified approach via fluid limit models. *Ann. Appl. Probab.*, 5(1):49–77, Feb. 1995.
- [7] Leonidas Georgiadis, Michael J Neely, and Leandros Tassiulas. *Resource allocation and cross-layer control in wireless networks*. Now Publishers Inc, 2006.

- [8] Swaroop Gopalam, Stephen V Hanly, and Philip Whiting. Distributed and local scheduling algorithms for mmwave integrated access and backhaul. *IEEE/ACM Trans. Netw.*, 30(4):1749–1764, Aug. 2022.
- [9] Manan Gupta, Anil Rao, Eugene Visotsky, Amitava Ghosh, and Jeffrey G Andrews. Learning link schedules in self-backhauled millimeter wave cellular networks. *IEEE Trans. Wireless Commun.*, 19(12):8024–8038, Dec. 2020.
- [10] Sean P Meyn and Richard L Tweedie. *Markov chains and stochastic stability*. Springer Science & Business Media, 2012.
- [11] Michael Neely. *Stochastic network optimization with application to communication and queueing systems*. Morgan & Claypool Publishers, 2010.
- [12] Michele Polese, Marco Giordani, Tommaso Zugno, Arnab Roy, Sanjay Goyal, Douglas Castor, and Michele Zorzi. Integrated access and backhaul in 5g mmwave networks: Potential and challenges. *IEEE Commun. Mag.*, 58(3):62–68, Mar. 2020.
- [13] Panagiotis Promponas, Tingjun Chen, and Leandros Tassiulas. On the optimization and stability of sectorized wireless networks. *IEEE Trans. Netw.*, Dec. 2025.
- [14] Ian P Roberts, Aditya Chopra, Thomas Novlan, Sriram Vishwanath, and Jeffrey G Andrews. Beamformed self-interference measurements at 28 ghz: Spatial insights and angular spread. *IEEE Trans. Wireless Commun.*, 21(11):9744–9760, Nov. 2022.
- [15] Ian P Roberts, Sriram Vishwanath, and Jeffrey G Andrews. Lonestar: Analog beamforming codebooks for full-duplex millimeter wave systems. *IEEE Trans. Wireless Commun.*, 22(9):5754–5769, Sep. 2023.
- [16] R Tyrrell Rockafellar. *Convex analysis*, volume 28. Princeton university press, 1997.
- [17] Yekaterina Sadovaya, Olga Vikhrova, Wei Mao, Shu-ping Yeh, Omid Semiari, Hosein Nikopour, Shilpa Talwar, and Sergey Andreev. Delay-aware link scheduling in iab

- networks with dynamic user demands. *IEEE Trans. veh. Technol.*, 73(10):15125–15139, Oct. 2024.
- [18] Mathew K Samimi, George R MacCartney, Shu Sun, and Theodore S Rappaport. 28 ghz millimeter-wave ultrawideband small-scale fading models in wireless channels. In *Proc. IEEE 83rd Veh. Technol. Conf. (VTC Spring)*, pages 1–6. IEEE, May 2016.
- [19] L. Tassiulas and A. Ephremides. Stability properties of constrained queueing systems and scheduling policies for maximum throughput in multihop radio networks. *IEEE Trans. Autom. Control*, 37(12):1936–1948, Dec. 1992.
- [20] Xianghao Yu, Jun Zhang, Martin Haenggi, and Khaled B. Letaief. Coverage analysis for millimeter wave networks: The impact of directional antenna arrays. *IEEE J. Sel. Areas Commun.*, 35(7):1498–1512, July 2017.
- [21] Dingwen Yuan, Hsuan-Yin Lin, Jörg Widmer, and Matthias Hollick. Optimal and approximation algorithms for joint routing and scheduling in millimeter-wave cellular networks. *IEEE/ACM Trans. Netw.*, 28(5):2188–2202, Oct. 2020.
- [22] Bibo Zhang and Ilario Filippini. Mobility-aware resource allocation for mmwave iab networks: A multi-agent reinforcement learning approach. *IEEE/ACM Trans. Netw.*, 32(4):3559–3574, Aug. 2024.
- [23] Junkai Zhang, Haifeng Luo, Navneet Garg, Abhijeet Bishnu, Mark Holm, and Tharmalingam Ratnarajah. Design and analysis of wideband in-band-full-duplex fr2-iab networks. *IEEE Trans. Wireless Commun.*, 21(6):4183–4196, June 2021.
- [24] Yongqiang Zhang, Mustafa A Kishk, and Mohamed-Slim Alouini. Freshness-aware energy efficiency optimization for integrated access and backhaul networks. *IEEE Trans. Wireless Commun.*, 23(10):14715–14728, Oct. 2024.