

Robust Downlink Scheduling for Low Latency Earth Observation in LEO Satellite Networks

Abstract—Low Earth Orbit satellites generate tens of terabytes of Earth observation data daily, all of which must be downlinked within strictly time-constrained contact windows. However, due to orbital geometry and ground station deployment constraints, satellite-ground contact windows are inherently short and intermittent, making the complete transmission of large volumes of data a fundamental challenge. Existing methods treat predicted link rates as deterministic inputs for downlink scheduling, ignoring skewed heavy-tailed errors that cause backlog accumulation and data loss under capacity overestimation. To address these challenges, this paper proposes a rolling horizon robust scheduling algorithm. We propose a conditional error distribution model that captures prediction error statistics and introduces risk-adjusted capacity to suppress overestimation with probabilistic guarantees. A rolling horizon framework then decomposes the stochastic dynamic programming problem into finite-window earliest-deadline-first greedy subproblems, achieving joint optimization of completion rate, latency, and throughput without prior knowledge of actual link rates. Experiments show that our algorithm achieves a 94.3% task completion rate, outperforming the state-of-the-art by 22.1%, with 73.4% higher throughput, 57.1% lower latency, and 75.3% less data loss.

Index Terms—LEO satellite networks, Integrated satellite and terrestrial networks, Downlink scheduling, Stochastic optimization

I. INTRODUCTION

Low Earth orbit (LEO) satellite networks have become critical infrastructure for global Earth observation (EO), with over 45% of satellites operating below 2,000 km and collecting tens of terabytes of remote sensing data daily [1], [2]. This data must be downlinked to ground-based EO data centers before distribution to over 4.53 billion users worldwide [3], as shown in Fig. 1. Many downstream applications impose strict timeliness requirements, where data delivered beyond task deadlines in scenarios such as disaster response, extreme weather forecasting, and maritime safety monitoring loses all operational value [4], [5]. However, LEO satellite-ground communication is intermittent, with

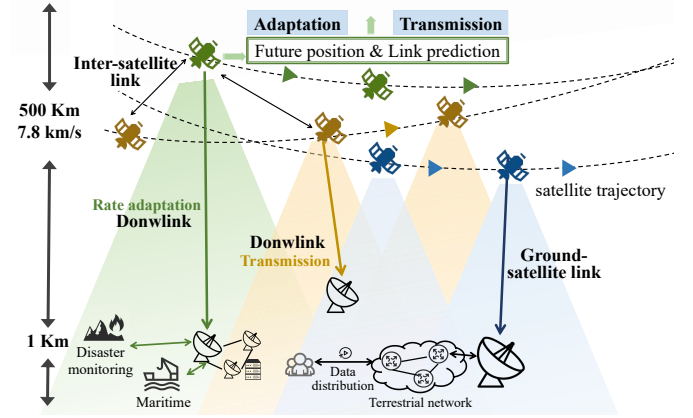


Fig. 1: Data delivery process in typical EO systems.

contact windows severely limited by orbital geometry and ground station deployment, making complete transmission of large-volume data within time constraints a fundamental challenge [6]. How to rationally schedule downlink priority and bandwidth allocation within these short contact windows is therefore a core problem.

To alleviate downlink bottlenecks, existing methods mainly follow two approaches. The first leverages on-board edge computing (OEC) to preprocess and compress data before transmission, yet as task volume grows and data requirements diversify, this approach cannot fundamentally eliminate the bottleneck [2], [7]. The second approach determines downlink scheduling decisions before the satellite-ground link communication window opens, using link prediction models built from elevation angles, historical SNR, and weather forecasts as scheduling basis [8]. However, actual link rates are inevitably affected by elevation variation, rain attenuation [9], [10], leading to deviations from predicted values, with our measured data showing rates falling significantly below predictions in more than 40% of windows. When the scheduler executes based on predicted rates, insufficient actual throughput causes data to accumulate in onboard buffers, increasing transmission latency and preventing tasks from completing within their deadlines. This mismatch between predicted and

actual link rates causes downlink scheduling failure in LEO constellations and is the core challenge.

Accordingly, we conduct an in-depth analysis of prediction error statistics and identify two properties overlooked. First, the error distribution exhibits skewness and heavy tails. Our measurements show that while prediction errors remain near 0 for most windows, they grow substantially under low elevation angles, rain attenuation, and multipath interference, forming a heavy tail that dominates data loss and excessive latency. Second, adjacent window prediction errors exhibit significant temporal correlation, driven by the continuity of orbital geometry and the persistence of weather conditions. Existing scheduling methods treat predicted link rates as exact quantities, implicitly assuming zero prediction error. This flaw is compounded by onboard data accumulation dynamics, where overestimated capacity leaves data undelivered, which is then either lost outright or carried over to subsequent windows. Consecutive overestimations from temporal correlation then let the backlog grow, causing multiple tasks to miss deadlines and triggering large-scale delivery failures.

To address the scheduling uncertainty introduced by link rate prediction errors, this paper proposes a rolling horizon robust scheduling algorithm (RH-EDF) that explicitly models the statistical structure of prediction errors and formalizes the scheduling problem as a stochastic dynamic programming problem. Specifically, we propose a conditional error distribution model for each satellite-ground link that dynamically updates the conditional quantile of the current window using the observed error from the most recent window, thereby transforming temporal error correlation into quantifiable risk information. To transform stochastic uncertainty into deterministic constraints, we propose risk-adjusted capacity, using a probability δ to ensure that scheduling decisions do not exceed the actual link capacity, thus suppressing data backlog caused by capacity overestimation. To overcome the computational intractability of infinite-horizon stochastic dynamic programming, we adopt a rolling horizon framework that decomposes the original problem into a finite-window deterministic earliest-deadline-first (EDF) subproblem. In each round, only the current window decision is executed, and error state is updated with actual observations, enabling the scheduler to continuously integrate the latest error information without true rates, achieving joint optimization of task completion rate, transmission latency, and throughput. Extensive experiments demonstrate that

RH-EDF achieves an average task completion rate of 94.3%, improving by 22.1% over the state-of-the-art baseline, with 73.4% higher throughput, 57.1% lower average latency, and 75.3% reduction in data loss rate.

In summary, the main contributions of our paper are:

- We identify two overlooked properties of link rate prediction errors, heavy-tailed skewness and inter-window temporal correlation, as the primary source of downlink scheduling uncertainty in satellite network.
- We propose a rolling horizon robust scheduling algorithm that models conditional error distributions, introduces risk-adjusted capacity, and decomposes the problem into tractable finite-window EDF subproblems.
- Extensive experiments show that RH-EDF outperforms state-of-the-art baselines in task completion rate, throughput, completion delay, and data loss rate.

II. RELATED WORK

The problem of downlink constraints on emerging LEO satellite constellations has been studied [7], [9], [11]. We review related work on downlink scheduling and selective downlinking and analyze their limitations.

A. Downlink scheduling

Downlink link scheduling has been addressed in [5], [12]–[16]. These systems either ignore time-varying link quality. AEROPATH [14] jointly optimizes satellite selection and inter-satellite routing under a ground-station-driven architecture to maximize throughput with low latency, yet assumes precise downlink capacity is known, disregarding fading and weather effects. Orbit-Cast [15] leverages large-scale constellations for low-latency data distribution via location-aware routing, but similarly relies on exact capacity predictions without modeling prediction errors. Tang et al. [16] propose multi-path cooperation to improve reliability, yet still lack capacity uncertainty modeling at the scheduling level. In summary, existing scheduling studies treat downlink capacity as a deterministic input and fail to adequately capture its time-varying nature.

B. Opportunistic communication and selective downlink

Opportunistic communication and selective downlink are important paradigms for coping with the intermittent connectivity of LEO satellites. Recently, [7] proposed offloading computation to satellites to reduce downlink load, such as pre-filtering building imagery before downlink. However, this approach conflicts with the Earth observation business model, where satellites

TABLE I: Related works

	Downlink scheduling	Capacity uncertainty	Multi-window cumulative	Safety margin
[14]	✓	×	×	×
[15]	✓	×	×	×
[7]	✓	×	×	×
[8]	✓	×	×	✓
[16]	✓	×	×	×
Ours	✓	✓	✓	✓

distribute raw data to users. Without knowledge of the end application, onboard pre-filtering risks discarding user-critical information. [8] introduced L2D2, a distributed hybrid ground station architecture using low-cost commodity hardware to achieve low-latency and robust downlink. L2D2 improves link quality prediction via a data-driven model accounting for multipath and occlusion [9]. Nevertheless, it discards prediction uncertainty in downlink decisions and absorbs errors in a fixed buffer, thus failing to handle systematic bias across consecutive windows or model the impact of prediction error on capacity decisions.

C. Summary

As summarized in Table 1, existing methods achieve efficient transmission with low latency and high throughput, yet their assumptions deviate from reality. L2D2 addresses both downlink scheduling and fixed safety margin, but it relies solely on link prediction without modeling the uncertainty and the multi-window cumulative effects of prediction errors on scheduling.

III. SYSTEM MODEL AND PROBLEM FORMULATION

A. Network model

Let the constellation consist of M LEO satellites, denoted by the set $\mathcal{S} = \{s_1, s_2, \dots, s_M\}$. The set of ground stations is denoted by $\mathcal{G} = \{g_1, g_2, \dots, g_N\}$, which represent distributed ground station nodes. Each satellite is equipped with an onboard buffer of finite capacity, with a maximum size denoted by $B^{\max}$, used for temporary storage of downlink datasets. The buffer state at time t is denoted by $B(t) \in [0, B^{\max}]$.

B. Intermittent communication model

For each satellite–ground station pair, contact windows form a non-overlapping sequence. The i -th win-

dow is defined as $w_i = (t_i^{\text{start}}, t_i^{\text{end}}, s, g)$, with duration $L_i = t_i^{\text{end}} - t_i^{\text{start}}$. The actual capacity of window i is

$$C_i = \int_{t_i^{\text{start}}}^{t_i^{\text{end}}} r_i(t) dt \approx \bar{r}_i \times L_i, \quad (1)$$

and the predicted capacity, derived from elevation angle, weather, and historical SNR, is

$$\hat{C}_i = \hat{r}_i \times L_i. \quad (2)$$

The deviation between $\hat{C}_i$ and C_i constitutes the root cause of scheduling uncertainty.

C. Link rate prediction error model

We define the prediction error ε_i as the relative deviation between the actual rate and the predicted rate,

$$r_i = \hat{r}_i \cdot (1 - \varepsilon_i), \quad (3)$$

$\varepsilon_i > 0$ indicates overestimation that actual rate lower than predicted, $\varepsilon_i < 0$ indicates underestimation that actual rate higher than predicted. ε_i is not symmetric, independent, and identically distributed noise, but it exhibits two characteristics.

Observation 1. Skewed heavy-tailed errors. Based on real-world satellite data collected from HEDE Aerospace, spanning 2023.01 to 2025.01 (63,072,000 records covering channel quality, satellite position), the distribution of ε_i exhibits a bimodal, skewed structure, with the main peak concentrated near $\varepsilon_i \approx 0$, indicating that the link quality of most windows matches the prediction. The heavy tail on the overestimation side ($\varepsilon_i > 0$) corresponds to low-elevation windows, precipitation attenuation (2–3 dB for X/Ka-band), and multipath obstruction. This tail risk, not Gaussian noise, dominates deadline violations, and the mean $\mathbb{E}[\varepsilon_i]$ fails to represent typical window behavior.

Observation 2. Temporal correlation of errors. Link rate errors of adjacent windows over the same ground station exhibit significant positive correlation, due to orbital geometry continuity (similar elevation profiles, low-elevation windows in clusters) and weather persistence (hourly-scale changes slower than window intervals). No such correlation exists for different ground stations. For each ground station g , the error sequence follows an autoregressive process,

$$\varepsilon_{g,j} = \mu_{\varepsilon,g} + \rho_g(\varepsilon_{g,j-1} - \mu_{\varepsilon,g}) + \eta_{g,j}. \quad (4)$$

where $\rho_g > 0$ captures the persistence of overestimation across consecutive windows. Consecutive overestimation is far more likely than under independence, amplifying systematic deadline violations.

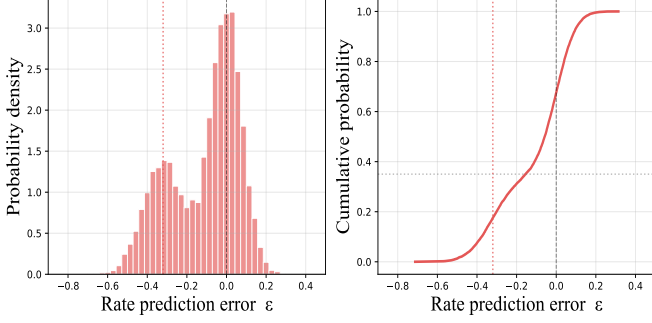


Fig. 2: Structural characteristics of link rate prediction.

D. Task-level queue and buffer model

Satellite s generates EO tasks $\mathcal{D}_s$, each described by a triplet $(t_d, \text{size}_d, \tau_d)$ denoting generation time, data size, and deadline. The on-board buffer maintains a task queue Q_i of incomplete tasks before window w_i , with scalar backlog $B_i = \sum_{d \in Q_i} \text{remain}_d$. At window w_i , newly generated data D_i is merged into the queue, and the transmission allocation x_i must satisfy

$$0 \leq x_i \leq \min(C_i, B_i + D_i). \quad (5)$$

The actual transmitted bytes are $\mathbb{T}_i = \min(x_i, C_i) = \min(x_i, \hat{C}_i(1 - \varepsilon_i))$. The backlog evolves as

$$B_{i+1} = \min(B^{\max}, \max(0, B_i + D_i - \mathbb{T}_i)), \quad (6)$$

where $\min(\cdot, B^{\max})$ ensures that the buffer does not exceed the on-board storage limit, with any excess treated as data loss, and $\max(0, \cdot)$ guarantees non-negativity.

E. Multi-window error accumulation model

When $x_i > C_i$ (i.e., the scheduler overestimates the capacity), we have $\mathbb{T}_i = C_i < x_i$, and the difference $x_i - C_i$ remains as backlog carried over to B_{i+1} , exerting pressure on subsequent windows. This recursive relationship reveals rollback accumulation. B_i does not reset automatically. When the scheduler overestimates a window's capacity ($\varepsilon_i > 0$), the undelivered data $x_i - \mathbb{T}_i > 0$ is carried over as backlog to the next window. If the scheduler remains unaware of this backlog when planning the next window (since it relies solely on the predicted value $\hat{C}_{i+1}$), the backlog continues to accumulate. When overestimation errors occur consecutively (Observation 2), the backlog grows beyond the compensation capacity of subsequent windows, ultimately causing tasks to miss their deadlines. Furthermore, once the backlog crosses a critical threshold, it causes multiple tasks to miss their deadlines simultaneously, rather than deteriorating gradually on a per-task basis.

F. Problem formulation

We formally cast the scheduling problem under rate prediction uncertainty. Let $B_i \in [0, B^{\max}]$ denote the buffer backlog (state) and $x_i \geq 0$ the transmission allocation (decision) at window w_i . The true capacity is

$$C_i = \hat{C}_i \cdot (1 - \varepsilon_i), \quad (7)$$

where $\hat{C}_i$ is the predicted capacity and $\varepsilon_i \sim F_\varepsilon$ is the prediction error with a skewed, temporally correlated distribution as characterized in §III-C.

Therefore, the scheduler aims to find an allocation strategy $\{x_i\}$ that simultaneously minimizes average delay, maximizes completion rate, and maximizes throughput under the link rate prediction error distribution F_ε . We formulate the optimization objective in the form of a normalized weighted sum,

$$\begin{aligned} \min_{\{x_i\}} \quad & \mathbb{E}_{F_\varepsilon} \left[w_1 \cdot \frac{\bar{\ell}_d}{\ell_{\text{ref}}} - w_2 \cdot \text{CR} - w_3 \cdot \frac{\text{TP}}{\text{TP}_{\max}} \right] \\ & \ell_d = t_d^{\text{finish}} - t_d, \quad \forall d \\ & \text{CR} = \frac{1}{|\mathcal{D}_s|} \sum_{d \in \mathcal{D}_s} \mathbf{1}[t_d^{\text{finish}} \leq \tau_d] \\ & \text{TP} = \frac{\sum_{i=1}^K \mathbb{T}_i}{\mathbb{T}_{\text{total}}} \\ \text{s.t.} \quad & \text{C1: } 0 \leq x_i \leq \min(C_i, B_i + D_i) \\ & \text{C2: } B_{i+1} = \min(B^{\max}, \\ & \quad \max(0, B_i + D_i - \min(x_i, C_i))) \\ & \text{C3: } 0 \leq B_i \leq B^{\max}, \quad \forall i \end{aligned} \quad (8)$$

where $\mathbb{T}_{\text{total}}$ is the total duration covered by all windows. ℓ_{ref} is the median of task deadlines, and $\text{TP}_{\max}$ is the theoretical maximum throughput, defined as the sum of true capacities of all windows divided by the total duration. w_1 , w_2 , and w_3 are weight coefficients, set equal by default. For each window w_i ($i = 1, 2, \dots, K$), the following constraints must hold. C1 is the data availability constraint. The allocated transmission volume cannot exceed the predicted capacity. C2 refers to buffer dynamics. The backlog in each window equals the remaining data from the previous window plus newly arrived data minus the actual transmitted data. Backlog follows the discrete-time queueing recurrence from §III-E with $B_1 = 0$, coupling scheduling decisions x_i , channel capacity C_i , and new arrivals D_i . C3 is the physical limit. The backlog must not exceed the onboard memory capacity.

IV. ALGORITHM DESIGN

To solve Problem Eq. 8, three major challenges must be addressed. (i) The true capacity C_i is unknown at scheduling time, so the constraint $x_i \leq C_i$ cannot be directly enforced. (ii) The objective is in expectation form and nonlinear, making an analytical solution infeasible. (iii) The state is coupled across windows since B_{i+1} depends on ε_i , preventing per-window independent optimization. We therefore transform the original stochastic problem into a tractable deterministic equivalent form, combined with rolling replanning.

A. Constraint transformation

Both constraints C1 and C2 involve the unknowable capacity C_i , and the objective contains expectations over ε_i , rendering a direct solution intractable. We therefore relax the hard constraint into a chance constraint, replacing the exact capacity uncertainty with a probabilistic guarantee that capacity overflow occurs with probability at most δ . In the original constraint C1, $B_i + D_i$ is the data availability known at the scheduling time, so it can be retained directly. The truly unknowable part is C_i . We replace the corresponding sub-constraint $x_i \leq C_i$ with a chance constraint,

$$P(x_i \leq \hat{C}_i(1 - \varepsilon_i)) \geq 1 - \delta. \quad (9)$$

Solving this chance constraint is equivalent to:

$$x_i \leq \hat{C}_i \cdot (1 - F_{\varepsilon_i}^{-1}(1 - \delta)) := \hat{C}_i^\delta, \quad (10)$$

where $F_{\varepsilon_i}^{-1}(1 - \delta)$ is the $(1 - \delta)$ -quantile of the ε_i distribution. Accordingly, C1 as a whole is transformed into a deterministic constraint:

$$C1': \quad 0 \leq x_i \leq \min(\hat{C}_i^\delta, B_i + D_i) \quad (11)$$

The setting of δ reflects the asymmetric cost of overestimation versus underestimation (Observation 1), with δ binned by window risk level $\delta = \{\delta_{\text{tight}}, \delta_{\text{loose}}\}$. δ_{tight} is high risk windows with low elevation angle, precipitation, and high historical error. And δ_{loose} is a low-risk window with a high elevation angle and clear weather. After this step, the unknowable capacity constraint in the original problem is replaced by a deterministic conservative capacity upper bound $\hat{C}_i^\delta$, and the problem becomes a solvable deterministic optimization.

Next, we address the determinization of the state transition. The actual transmission amount $\mathbb{T}_i = \min(x_i, C_i)$ in C2 similarly contains the random C_i . Since C1 guarantees $x_i \leq \hat{C}_i^\delta$, and $\hat{C}_i^\delta$ is a conservative value set at confidence level $1 - \delta$, we have $x_i \leq \hat{C}_i^\delta \leq C_i$ with

probability at least $1 - \delta$, so $\mathbb{T}_i \approx x_i$ holds at the same confidence level. Accordingly, C2 is simplified to the deterministic recurrence,

$$C2': B_{i+1} = \max(0, B_i + D_i - \min(\hat{C}_i^\delta, B_i + D_i)) \quad (12)$$

Unlike the physical buffer model, we no longer truncate and discard excess data at $B^{\max}$. Instead, overflow is treated as a penalty in the objective. This avoids corrupting task completion determination, as physical data loss would invalidate the integrity check of each task.

After the above substitutions, B_i becomes deterministic given x_i , independent of ε_i . All three performance metrics are computed from deterministic queue simulation, so the expectation operator is eliminated and the problem reduces to a deterministic optimization.

$$\min_{x_i} \quad w_1 \cdot \frac{\bar{\ell}_d}{\ell_{\text{ref}}} - w_2 \cdot \text{CR} - w_3 \cdot \frac{\text{TP}}{\text{TP}_{\max}} \quad (13)$$

$$\text{s.t.} \quad C1', C2', \quad 0 \leq B_i \leq B^{\max} \quad \forall i, \quad B_1 = 0 \quad (14)$$

B. Link rate prediction error distribution estimation

Section IV-A introduces the capacity reduction formula $\hat{C}_i^\delta = \hat{C}_i - F_{\varepsilon_i}^{-1}(1 - \delta)$, where the key term is the quantile of ε_i . We convert this quantile into a computable value. Since the temporal correlation of errors is station-specific rather than global, we perform estimation per ground station. For each link of ground station g , we calculate the mean $\mu_{\varepsilon,g}$, standard deviation $\sigma_{\varepsilon,g}$, and autocorrelation ρ_g from its historical error sequence. Based on the autoregressive model, we construct the conditional distribution $F_{\varepsilon_i|\varepsilon_{g,\text{prev}}^{\text{obs}}}$ with Eq. (4). The quantile term becomes $\hat{C}_i^\delta = \hat{C}_i - F_{\varepsilon_i|\varepsilon_{g,\text{prev}}^{\text{obs}}}^{-1}(1 - \delta)$, where making $\hat{C}_i^\delta$ directly computable. During online operation, each window's observed error $\varepsilon_i^{\text{obs}} = 1 - C_i^{\text{obs}}/\hat{C}_i$ updates the station's conditional state $\varepsilon_{g,\text{prev}}^{\text{obs}}$ for future visits.

C. Rolling horizon scheduling algorithm

After constraint transformation, each satellite faces a finite-horizon deterministic task scheduling problem. By the single-link access constraint, the scheduling of M satellites is mutually independent and can be solved separately. Algorithm 1 outlines the overall rolling horizon scheduling framework. At each window, it updates the link prediction error state using telemetry from the most recent pass, computes risk-adjusted capacities $\hat{C}_k^\delta$ for the next H windows via the conditional quantile $F^{-1}(1 - \delta)$, and invokes Algorithm 2 to solve a deterministic H -window subproblem. Only the first window's decision is executed, while the remaining windows serve as a

Algorithm 1 Rolling horizon scheduling algorithm(RH)

Input: $\mathcal{S}$, G , task sets $\mathcal{D}_s$, horizon length H , risk tolerance δ , error distribution parameters $\{\mu_{\varepsilon,g}, \sigma_{\varepsilon,g}, \rho_g\}$

Output: Transmission schedule and completion times

- 1: **for** each satellite $s \in \{1, \dots, M\}$ in parallel **do**
- 2: Initialize $Q_{s,1} \leftarrow \emptyset$, $B_{s,1} \leftarrow 0$
- 3: **for** $i = 1, 2, \dots, K_s$ **do**
- 4: $g_i \leftarrow$ ground station for window $w_{s,i}$
- 5: **if** g_i has been visited before **then**
- 6: $\varepsilon_{g_i, \text{prev}}^{\text{obs}} \leftarrow 1 - C_{g_i, \text{prev}}^{\text{obs}} / \hat{C}_{g_i, \text{prev}}^{\delta}$
- 7: **end if**
- 8: **for** $k = i$ to $\min(i + H - 1, K_s)$ **do**
- 9: $\hat{C}_k^{\delta} \leftarrow \hat{C}_i \cdot (1 - q_k^{\delta})$ where $q_k^{\delta} = F_{\varepsilon|o}^{-1}(1 - \delta_k)$
- 10: **end for**
- 11: Add arriving tasks $d \in \mathcal{D}_s$ with $t_d \leq w_{s,i}$
- 12: start to $Q_{s,i}$
- 13: $(x_i^*, \dots) \leftarrow \text{TQA}(Q_{s,i}, B_{s,i}, \{\hat{C}_k^{\delta}\}, H)$
- 14: Execute transmission: $x_i^{\text{act}} \leftarrow \min(x_i^*, C_{g_i}, B_{s,i} + D_i)$
- 15: Update $Q_{s,i+1}$, record t_d^{finish} for completed tasks
- 16: **end for**
- 17: **end for**

lookahead buffer. After execution, the actual transmission is capped by the realized capacity C_i , and the queue is updated with any unfinished data carried forward.

RH innovates in three aspects. First, it maintains per-station conditional error distributions that exploit temporal correlation across windows rather than treating errors independently. Second, risk-adjusted capacity $\hat{C}_i^{\delta}$ accounts for heavy-tailed skewness, bounding overflow probability to at most δ . Third, the rolling horizon with lookahead couples short-term execution with future planning, bridging stochastic uncertainty and deterministic optimization without sacrificing statistical fidelity.

D. Subproblem solving: task-level queue allocation

The subproblem solved by task-level queue allocation (TQA) at w_i is defined as follows. Given the queue Q (containing backlogged tasks and newly arrived ones up to $w_i.\text{start}$), the set of risk-adjusted capacities $\{\hat{C}_k^{\delta}\}$ for the next H windows $k = i, \dots, i + H - 1$, and the deadlines τ_d of all tasks, the solver determines the allocation x_k^* for each window k within the horizon. The allocation follows an earliest deadline first (EDF) greedy policy. All tasks are first sorted in ascending order of τ_d , with ties broken by remaining bytes in ascending order. For each task, the algorithm greedily assigns data

Algorithm 2 EDF-based Subproblem Greedy Algorithm

Input: $Q_{s,i}$, $\hat{C}_k^{\delta}$, H

Output: Optimal allocation $\{x_k^*\}$

- 1: Sort tasks by τ_d ascending
- 2: **for** each task d in sorted order **do**
- 3: **for** $k = i$ to $\min(i + H - 1, K_s)$ **do**
- 4: **if** $w_{s,k}.\text{start} \geq t_d$ **then**
- 5: Allocate $\min(\text{remain}_d, \hat{C}_k^{\delta, \text{rem}})$ to window k
- 6: **if** task d complete **then**
- 7: Record t_d^{finish} ; **break**
- 8: **end if**
- 9: **end for**
- 10: **end for**
- 11: **end for**

to the earliest feasible window after t_d with remaining capacity, marking the task complete once fully allocated or pending otherwise. Output is per-window allocation vector $\{x_k^*\}$, from which only x_i^* is executed.

V. PERFORMANCE EVALUATION

A. Experimental setting

We use SpaceX Starlink Phase I (1584 satellites, 72 orbits, 550 km) as the primary LEO constellation, with remote sensing satellite orbits [15] and 173 ground stations [17]. EO traffic is generated from NASA LANCE datasets [18], [19] for emergency response and disaster management. Task sizes range from 50 to 500 MB, with deadlines split into urgent (30%, 1–3h), normal (50%, 4–8h), and loose (20%, 12–24h). For comparison, we implement three benchmark algorithms. The distributed ground station (DGS) [8] employs a hybrid ground station model integrated with a link prediction mechanism to reduce downlink latency. OrbitCast [15] leverages inter-satellite links to forward data to other satellites that have upcoming downlink windows when predicted transmission rate is insufficient to complete data delivery within current window. Oracle uses actual capacity C_i with the same EDF-based greedy allocation, serving as the upper bound for completion rate and throughput and the lower bound for delay. We evaluate three metrics: completion rate (CR), system throughput (TP), and latency, as defined in Eq. (8).

Link rate estimation uses a link budget model. The predicted rate $\hat{r}_i$ is derived from ground station parameters and window elevation angle, which determine slant range d_i , yielding SNR via free-space path loss and then $\hat{r}_i$ via the Shannon formula. The actual rate

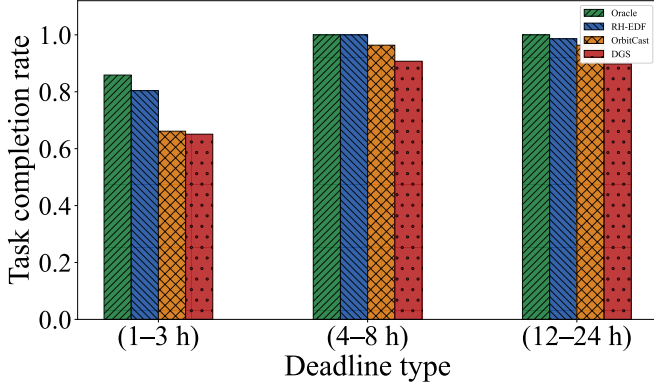


Fig. 3: Task completion rate under different deadline tiers.

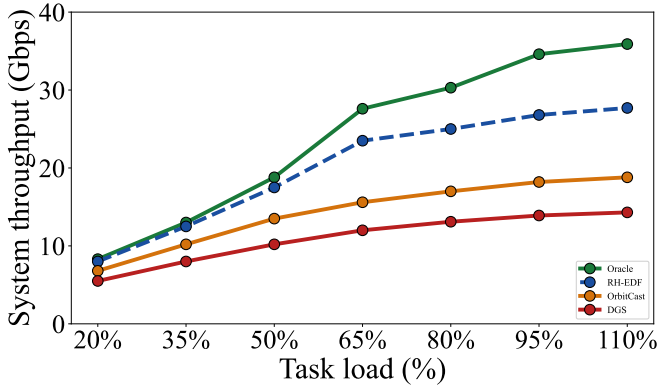


Fig. 4: System throughput under varying task loads.

superimposes structured errors onto $\hat{r}_i$: a base prediction error following a skewed Beta distribution (mean 0.06, 90th percentile 0.41), additional attenuation from Uniform(0.1, 0.3) at low elevation angles ($<15^\circ$), and weather attenuation from Uniform(0.05, 0.15) triggered with probability 0.15. All parameters are calibrated from the L2D2 measurement report [8].

B. Results and analysis

Analysis of task completion rate under different deadline tiers. Fig. 3 shows that RH-EDF consistently outperforms the baselines across all deadline tiers and closely tracks the Oracle. Under tight deadlines, where compensation windows are scarce, the proposed coordinated optimization mechanism, which combines outer-tier risk-adjusted capacity $\hat{C}_k^\delta$ to cap allocation and inner-tier EDF ordering with H-window compensation, enables data delivery within the shortest time. In addition, the probabilistic safety margin ensures 0 data loss.

System throughput analysis. This experiment sweeps task load from 20% to 110%. Under low-to-moderate loads, all algorithms differ by less than 1% as predicted capacity is not yet strained. Beyond 65%, prediction error skewness triggers capacity truncation: DGS and OrbitCast lack risk-adjusted capacity $\hat{C}_k^\delta$ and suffer exponentially increasing truncation events and throughput collapse, while RH-EDF keeps allocations below the quantile threshold, making truncation negligible. In the 65–80% range, RH-EDF delivers 95.8% more throughput than DGS and 47–51% more than OrbitCast. Under high load, EDF-ordered allocation directs available capacity to the tightest-deadline tasks, preventing persistent backlog accumulation.

Analysis of completion delay under varying ground station densities. Fig. 5 shows that RH-EDF achieves significantly lower delay than the baselines. This advantage stems from the error-state update in Algorithm 1, which identifies overestimation windows in advance, combined with EDF ordering and H-window compensation that prevent queuing delays caused by over-allocation in DGS and OrbitCast. As ground station density increases, visible windows become more frequent and all algorithms exhibit lower delays. However, RH-EDF prioritizes each new window capacity for the most urgent tasks, achieving the largest delay reduction of 54.3% as the number of ground stations increase.

Data loss rate under varying task loads. Fig. 6 shows that RH-EDF maintains data loss rate below 0.05 across all loads, while baselines exceed 0.10 beyond 30% load. Baselines allocate based on predicted capacity, causing backlog to accumulate under consecutive overestimation and eventually overflow the buffer. RH-EDF suppresses over-allocation at the source via risk-adjusted capacity $\hat{C}_k^\delta$, while EDF ordering prioritizes urgent tasks, keeping backlog below the buffer limit across all load conditions.

VI. CONCLUSION

We propose RH-EDF, a rolling horizon robust scheduling algorithm that addresses link rate prediction uncertainty in LEO constellation downlink scheduling. By modeling conditional error distributions and introducing risk-adjusted capacity, the stochastic scheduling problem is decomposed into tractable finite-window subproblems. Experiments show that RH-EDF improves task completion rate by 22.1% and reduces average delivery latency by 57.1% over baselines, demonstrating robust performance under uncertain link conditions.

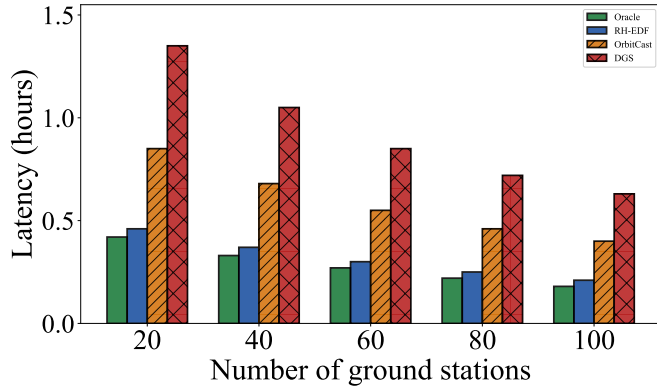


Fig. 5: Average completion delay under varying ground station densities.

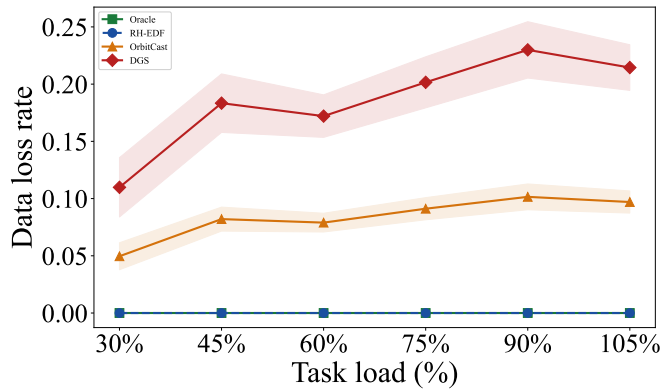


Fig. 6: Data loss rate under varying task loads.

and laying the foundation for future work on adaptive scheduling in LEO constellations.

REFERENCES

- [1] C. Yang, Q. Zhang, Q. Sun, S. Ouyang, A. Zhou, S. Wang, and M. Xu, "Fedcl+: Tackling onboard label constraints for accurate federated satellite computing," *IEEE Transactions on Services Computing*, 2025.
- [2] X. Yang, J. Feng, L. Liu, Q. Pei, S. Mumtaz, K. Li, and S. Dustdar, "Irs-assisted hyperspectral image processing in satellite edge computing services," *IEEE Transactions on Services Computing*, 2025.
- [3] Y. Li, Y. Chen, J. Yang, J. Zhang, B. Sun, L. Liu, H. Li, J. Wu, Z. Lai, Q. Wu, *et al.*, "Small-scale leo satellite networking for global-scale demands," in *Proceedings of the ACM SIGCOMM 2025 Conference*, pp. 917–931, 2025.
- [4] P. Deng, X. Gong, and X. Que, "A mobility adaptive service placement scheme in satellite edge computing network," in *2022 IEEE 24th Int Conf on High Performance Computing Communications; 8th Int Conf on Data Science Systems; 20th Int Conf on Smart City; 8th Int Conf on Dependability in Sensor, Cloud Big Data Systems Application (HPCC/DSS/SmartCity/DependSys)*, pp. 1430–1435, 2022.
- [5] C. Wang, Y. Peng, J. Wu, L. Liu, S. Mumtaz, M. Dong, and M. Guizani, "Deterministic scheduling and reliable routing for smart ocean services in maritime internet of things: A cross-layer approach," *IEEE Transactions on Services Computing*, 2024.
- [6] F. Tang, X. Li, L. Chen, J. Liu, M. Gao, Y. Yang, W. Xu, and H. Zhang, "Reinforcement learning based control domain division in leo satellite networks," in *2023 IEEE International Conference on High Performance Computing Communications, Data Science Systems, Smart City Dependability in Sensor, Cloud Big Data Systems Application (HPCC/DSS/SmartCity/DependSys)*, pp. 303–310, 2023.
- [7] B. Denby and B. Lucia, "Orbital edge computing: Nanosatellite constellations as a new class of computer system," in *Proceedings of the Twenty-Fifth International Conference on Architectural Support for Programming Languages and Operating Systems*, pp. 939–954, 2020.
- [8] D. Vasisht, J. Shenoy, and R. Chandra, "L2d2: Low latency distributed downlink for leo satellites," in *Proceedings of the 2021 ACM SIGCOMM 2021 Conference*, pp. 151–164, 2021.
- [9] K. Devaraj, M. Ligon, and E. Blossom, "Planet high speed radio: Crossing gbps from a 3u cubesat," 2019.
- [10] Y. Li, H. Li, L. Liu, W. Zhao, Y. Chen, J. Wu, Q. Wu, J. Liu, Z. Lai, *et al.*, "A networking perspective on starlink's self-driving leo mega-constellation," in *Proceedings of the 29th Annual International Conference on Mobile Computing and Networking*, pp. 1–16, 2023.
- [11] H. Li, Y. Li, Y. Liu, B. Deng, Y. Li, X. Li, and S. Zhao, "Earth observation satellite downlink scheduling with satellite-ground optical communication links," *IEEE Transactions on Aerospace and Electronic Systems*, vol. 61, no. 2, pp. 2281–2294, 2025.
- [12] R. A. d. Carvalho, "Optimizing the communication capacity of a ground station network," *Journal of Aerospace Technology and Management*, vol. 11, p. e2319, 2019.
- [13] D. Karapetyan, S. M. Minic, K. T. Malladi, and A. P. Punnen, "Satellite downlink scheduling problem: A case study," *Omega*, vol. 53, pp. 115–123, 2015.
- [14] W. Liu, Q. Wu, Z. Lai, H. Li, Y. Li, and J. Liu, "Enabling ubiquitous and efficient data delivery by leo satellites and ground station networks," in *GLOBECOM 2022-2022 IEEE Global Communications Conference*, pp. 687–692, IEEE, 2022.
- [15] Z. Lai, Q. Wu, H. Li, M. Lv, and J. Wu, "Orbitcast: Exploiting mega-constellations for low-latency earth observation," in *2021 IEEE 29th International Conference on Network Protocols (ICNP)*, pp. 1–12, IEEE, 2021.
- [16] F. Tang, H. Zhang, and L. T. Yang, "Multipath cooperative routing with efficient acknowledgement for leo satellite networks," *IEEE Transactions on Mobile Computing*, vol. 18, no. 1, pp. 179–192, 2018.
- [17] "Satnogs: Open source global network of satellite ground-stations." <https://satnogs.org/>. Accessed: 2026-07-09.
- [18] "Misr near real-time (nrt) data." <https://earthdata.nasa.gov/earth-observation-data/near-realtime/download-nrt-data/misr-nrt>. Accessed: 2026-07-09.
- [19] "Modis near real-time (nrt) data." <https://earthdata.nasa.gov/earth-observation-data/near-realtime/download-nrt-data/modis-nrt>. Accessed: 2026-07-09.