

Multimodal Remote Inference

Keyuan Zhang*, Yin Sun[†], Bo Ji*

*Dept. of Computer Science, Virginia Tech [†]Dept. of Electrical and Computer Engineering, Auburn University
Email: {keyuanz, boji}@vt.edu, yzs0078@auburn.edu

Abstract—We consider a remote inference system with multiple modalities, where a multimodal machine learning (ML) model performs real-time inference using features collected from remote sensors. When sensor observations evolve dynamically over time, fresh features are critical for inference tasks. However, timely delivery of features from all modalities is often infeasible because of limited network resources. Towards this end, in this paper, we study a two-modality scheduling problem that seeks to minimize the ML model’s inference error, expressed as a penalty function of the Age of Information (AoI) vector of the two modalities. We develop an index-based threshold policy and prove its optimality. Specifically, the scheduler switches to the other modality once the current modality’s index function exceeds a predetermined threshold. We show that both modalities share the same threshold and that the index functions and the threshold can be computed efficiently. Our optimality results hold for general AoI functions (which could be non-monotonic and non-separable) and heterogeneous transmission times across modalities. To demonstrate the importance of considering a task-oriented AoI function, we conduct numerical experiments based on robot state prediction and compare our policy with round-robin and uniform random policies (both are oblivious to the AoI and the inference error). The results show that our policy reduces inference error by up to 55% compared with these baselines.

Index Terms—Multimodal Remote Inference; Age of Information; Scheduling

I. INTRODUCTION

The advent of sixth-generation (6G) technology, along with advances in artificial intelligence (AI), enables *remote inference* in various intelligent applications, such as autonomous transportation, unmanned mobility, and industrial automation [1]. As illustrated in Fig. 1, a remote inference system uses AI models for inference tasks (e.g., monitoring, reasoning, and decision-making) based on features transmitted from remote sensors. For instance, traffic prediction relies on near real-time forecasts of traffic status (e.g., speed, flow, and demand) based on spatio-temporal road data [2]. When the sensor observations change dynamically over time, timely data delivery is critical. For example, autonomous driving requires timely updates on vehicle positions and velocities to ensure safety; healthcare monitoring relies on real-time vital signs for timely alerts.

Moreover, such complex tasks often involve multiple data modalities, such as audio, visual, 3D points (e.g., LiDAR or RADAR), and motion (e.g., IMU) data. Each modality provides complementary information to enhance the overall inference accuracy. Take object detection and tracking as an example: while color images (RGB) capture the shape and

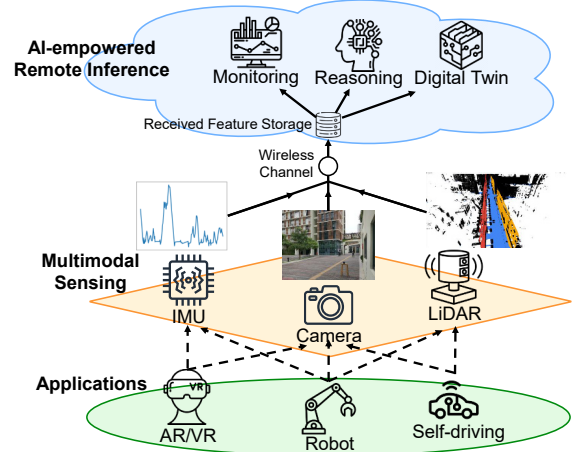


Fig. 1: A multimodal remote inference system

appearance of objects, LiDAR images offer depth information [3]. To fully exploit information from multiple modalities, machine learning (ML) based fusion techniques have been widely adopted, as they can effectively extract modality-specific information and capture cross-modal correlations using architectures such as deep neural networks [4].

Despite the advantages of multimodal ML, its application in remote inference systems remains limited in practice. One primary challenge is limited network resources (e.g., bandwidth), which makes it hard for multimodal ML models to access fresh features across all modalities simultaneously. Therefore, it is important to schedule transmissions for the modality that contributes most to the ML model’s performance. Another challenge is the lack of an accurate model that captures the relationship between feature freshness and inference accuracy [5]. That is, it is unclear which feature, at which time, is most *valuable* to the model. One main reason is the complexity of multimodal ML models, especially deep neural networks.

To address these challenges, prior studies on remote inference have offered key insights into how information freshness affects ML performance [1], [6], [7]. In these studies, *Age of Information (AoI)*, introduced in [8], serves as a crucial metric for measuring freshness. A key finding from these works reveals that in supervised learning, ML inference error can be expressed as a (possibly non-monotonic) function of AoI. This observation enables us to shift the goal from directly optimizing ML inference error to minimizing an AoI-based penalty function (which could be non-monotonic). Moreover,

the function depends on the AoI of each modality and may be non-additive (i.e., not a weighted sum of single-modality AoI functions); see Section V-B for an example. The generality of this function further complicates the scheduling problem.

To that end, in this paper, we study a key research question: *How to effectively schedule transmissions for multiple modalities to minimize the ML model's inference error under the network bandwidth constraint?* We answer this question by making the following main contributions:

- We formulate a two-modality scheduling problem that seeks to minimize inference error, which is a general function of the AoI of both modalities. The two-modality setting (e.g., visual-audio or RGB-LiDAR) is common in ML, but remains *underexplored* in remote inference.
- We develop an optimal *index-based threshold policy*. Specifically, one modality is scheduled continuously until its index function exceeds a predetermined threshold, and then the scheduler switches to the other modality. Interestingly, the two modalities share the same threshold. Both the index functions and the threshold can be computed efficiently. Our results hold for general AoI functions (which could be *non-monotonic* and *non-separable*) and *heterogeneous* transmission times across modalities.
- We conduct numerical experiments based on robot state prediction to demonstrate the importance of optimizing a general function of AoI. The results show that our policy reduces inference error by up to 55% compared with round-robin and uniform random policies (both are agnostic to the AoI and the inference error). We believe that these results provide valuable insights into improving the accuracy of multimodal remote inference through the optimization of task-oriented AoI functions.

II. RELATED WORK

AoI penalty functions. *Age of Information (AoI)* is a widely adopted metric for quantifying information freshness (see [1] and references therein, as well as a comprehensive survey [9]). Nevertheless, the relationship between information *freshness* and its *value* to the application is still not well understood. To that end, in [10], Sun and Cyr suggested several non-decreasing, non-linear AoI functions to capture the value of fresh data in various applications. Researchers have also proposed various metrics, such as Age of Incorrect Information (AoII) [11], to quantify the value of information in different systems (see [12] for a comprehensive survey).

More recent research has explored the impact of freshness in remote inference systems. In the seminal work [6], [7], Shisher *et al.* demonstrated that inference error can be expressed as a general AoI penalty function, which may not be monotonic. Based on this observation, the authors studied a transmission scheduling problem, aiming to minimize a general AoI penalty function. In [13], the joint optimization of feature length selection and transmission scheduling was further studied. These prior works assume that each task corresponds to a single source, which can be viewed as a *single-modality* setting. Optimizing a general AoI function for multiple modalities

remains underexplored. We take a step further by considering a general AoI function of two modalities, which is non-separable and thus adds an extra layer of complexity to the problem.

Multi-source AoI optimization. There have also been efforts to optimize AoI in multi-source systems, such as queueing systems [14], broadcast networks [15], and remote inference systems [6], [7], [13]. In [14], Sun *et al.* studied age-optimal online scheduling in a multi-flow, multi-server queueing system, where the age function is time-dependent, symmetric, and non-decreasing. In [15], Kadota *et al.* showed that the Maximum Age First (MAF) policy is optimal in a homogeneous network, and a suboptimal Whittle's index method was further applied for heterogeneous scenarios. In [13], Shisher *et al.* proposed the Net Gain policy for jointly optimizing feature length selection and source scheduling in multi-source scenarios. Although source and modality scheduling share similarities, our work considers a more general age function. Existing works on multi-source scheduling consider an additive or a non-decreasing age function, whereas a non-monotonic and non-additive age function naturally arises in the multimodal setting we consider. Therefore, approaches such as Whittle's index and Net Gain policy, which rely on the additive assumption, may not be directly applicable to our problem.

Additionally, scheduling for correlated sources is related to our work. Prior studies in correlated sources often model the correlation explicitly (e.g., one source may contain another's information) [16], [17]. In contrast, in our setting, the correlation among modalities is implicitly captured by the AoI-based penalty function, thus requiring a new algorithmic design.

III. SYSTEM MODEL AND PROBLEM FORMULATION

In this section, we model the multimodal remote inference system and formulate a two-modality scheduling problem.

As depicted in Fig. 2, we consider a multimodal remote inference system, where two sensors send features to a receiver over a wireless channel. Each sensor corresponds to a distinct modality $m \in \{1, 2\}$. Let time be slotted, indexed by $t = 0, 1, \dots$. At each time t , each modality m generates one feature $X_{m,t}$ from a given feature set $\mathcal{X}_m$; the joint feature set is defined as $\mathcal{X} := \mathcal{X}_1 \times \mathcal{X}_2$. Meanwhile, on the receiver side, a predictor (e.g., a neural network) uses the *freshest received* features from both modalities to infer the current target value Y_t from a given target set $\mathcal{Y}$. Because of bandwidth constraints, the scheduler can only select one modality for transmission at any given time. Suppose the n -th transmission starts at time S_n and completes at time D_n . At time S_n , the scheduler selects only one modality for transmission, and the decision is denoted by $a_n \in \{1, 2\}$. We assume that the scheduler always transmits the freshest feature, i.e., X_{a_n, S_n} , from modality a_n .

Moreover, let T_m denote the *fixed* transmission time for modality m , which is a positive finite integer. Since the feature size varies, T_1 and T_2 may differ. We assume a reliable channel and do not consider preemption during transmission. That is, the receiver successfully receives the n -th selected feature from modality a_n at time D_n , and the transmission duration is T_{a_n} . We also assume a work-conserving system: when the

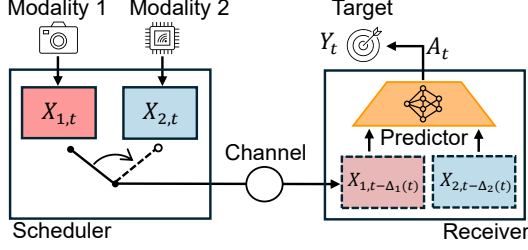


Fig. 2: System model

current transmission completes, the next transmission begins immediately, i.e., $S_{n+1} = D_n$ for every n .

We use *age of information (AoI)* to quantify information freshness, defined as the time elapsed since the freshest received feature was generated [18]. For each modality m , we denote its AoI at the receiver at time t as $\Delta_m(t) \in \mathbb{Z}_+$, where $\mathbb{Z}_+$ denotes the set of positive integers. According to the definition, the AoI of modality m at the receiver resets to its transmission time T_m whenever the receiver receives a feature from modality m (i.e., $a_n = m$ and $t = D_n$ for some n); otherwise, the AoI increases by 1. That is, the AoI of each modality m evolves as

$$\Delta_m(t) = \begin{cases} T_m & \text{if } a_n = m \text{ and } t = D_n, \\ \Delta_m(t-1) + 1 & \text{otherwise.} \end{cases} \quad (1)$$

The AoI vector of two modalities, denoted by $\Delta(t)$, is then defined as $(\Delta_1(t), \Delta_2(t)) \in \mathbb{Z}_+^2$.

Inference error. Define

$$\mathbf{X}_{t-\Delta(t)} := (X_{1,t-\Delta_1(t)}, X_{2,t-\Delta_2(t)}).$$

In order to predict the target $Y_t \in \mathcal{Y}$, the predictor $\phi: \mathcal{X} \times \mathbb{Z}_+^2 \mapsto \mathcal{A}$ outputs an action A_t from a given action set $\mathcal{A}$; the action is determined based on the freshest received features $\mathbf{X}_{t-\Delta(t)} \in \mathcal{X}$ and the associated AoI vector $\Delta(t) \in \mathbb{Z}_+^2$. The performance of the prediction is evaluated using a loss function $\ell(y, a)$, which quantifies the inference error incurred if the predictor selects action $a \in \mathcal{A}$ while the true target value is $y \in \mathcal{Y}$. For example, the action can be a probability distribution Q_Y in the space $\mathcal{Y}$, with the associated logarithmic loss function $\ell_{\log}(y, Q_Y) := -\log Q_Y(y)$. Alternatively, the action can be an estimate $\hat{y} \in \mathcal{Y}$ of the true target value $y \in \mathcal{Y}$, with the associated quadratic loss function $\ell_2(y, \hat{y}) := (y - \hat{y})^2$.

We assume that the process $\{(Y_t, X_t), t = 0, 1, \dots\}$ is *stationary*. This implies that the inference error is time-invariant. Second, the processes $\{(Y_t, X_t), t = 0, 1, \dots\}$ and $\{\Delta(t), t = 0, 1, \dots\}$ are *independent*. This holds when the scheduling policy does not know the feature or the target value (i.e., signal-agnostic). Under these two assumptions, the expected inference error at time t can be expressed as a function of the AoI vector [7]. We use $L: \mathbb{Z}_+^2 \mapsto \mathbb{R}$ to denote the expected inference error function, where $\mathbb{R}$ is the set of real numbers. For every AoI vector δ , function $L(\delta)$ is defined as

$$L(\delta) := \mathbb{E}_{Y, \mathbf{X} \sim \mathbb{P}(Y_t, \mathbf{X}_{t-\Delta(t)})} [\ell(Y, \phi(\mathbf{X}, \delta))],$$

where $\mathbb{P}(Y_t, \mathbf{X}_{t-\Delta(t)})$ denotes the joint stationary distribution of the target and the feature used for inference. The function L can be quite general; we only require it to be uniformly bounded, as stated below:

Assumption 1. *There exists a finite constant M such that $|L(\delta)| \leq M$ for all AoI vectors $\delta \in \mathbb{Z}_+^2$.*

Assumption 1 ensures the existence of an optimal policy in our analysis; this is also practical as ML models commonly apply preprocessing techniques to keep loss bounded.

Scheduling policy. The policy π is represented as a sequence of modality choices, i.e., $\pi := (a_0, a_1, a_2, \dots)$. We focus on scheduling policies that satisfy three conditions. (i) *Signal-agnostic*: The scheduler does not know the signal value (i.e., the feature or the target value) at any time. This is practical when these signals are private to the scheduler. (ii) *Causal*: Each decision a_n relies solely on the current and historical AoI (i.e., $\Delta(0), \Delta(1), \dots, \Delta(S_n)$) (iii) The scheduler knows the expected inference error function. In practice, this function can be numerically estimated by calculating the average loss over the dataset for each AoI vector. Let Π denote the set of all policies satisfying these conditions.

Objective. We aim to design a scheduling policy in Π that minimizes the time-averaged expected inference error over an infinite horizon. We define this problem as Problem **OPT**:

$$\mathbf{OPT} \quad \bar{L}_{\text{opt}} := \inf_{\pi \in \Pi} \limsup_{T \rightarrow \infty} \frac{1}{T} \mathbb{E}_{\pi} \left[\sum_{t=0}^{T-1} L(\Delta(t)) \right].$$

IV. OPTIMAL SCHEDULING POLICY DESIGN

The roadmap for this section is as follows. We first reformulate the problem and cast it as a Semi-Markov Decision Process (SMDP), as the reformulated problem is simpler to solve. Then, we present our policy, prove its optimality, and offer interpretations of our policy.

A. Problem Reformulation

At first glance, the problem requires selecting a modality at each delivery time. We next show that it can be reduced to two decisions, each associated with a special state.

Consider the following Semi-Markov Decision Process (SMDP). It has two states $\Delta_{1,\text{re}}, \Delta_{2,\text{re}} \in \mathbb{Z}_+$, defined as $\Delta_{1,\text{re}} := (T_1, T_1 + T_2)$ and $\Delta_{2,\text{re}} := (T_1 + T_2, T_2)$. Define m' as the complementary modality of modality m , i.e., $m' \in \{1, 2\} \setminus \{m\}$. Whenever the system reaches the AoI state $\Delta_{m,\text{re}}$ for some m , the scheduler selects $\tau \in \mathbb{N}$, indicating the number of consecutive transmissions of modality m before switching to modality m' . That is, the scheduler selects modality m for τ consecutive transmissions, followed by one transmission for modality m' . It turns out that the process always *restarts* at state $\Delta_{m',\text{re}}$, regardless of the value of the decision τ . This is because, according to the AoI evolution in Eq. (1), upon the delivery of the transmission from modality m' , the AoI of modality m' always resets to $T_{m'}$, while the AoI of modality m always increases from T_m to $T_m + T_{m'}$. Hence, we refer to the states $\Delta_{1,\text{re}}, \Delta_{2,\text{re}}$ as the *restart states*.

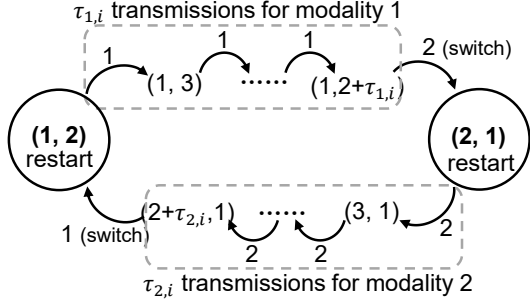


Fig. 3: AoI transition during the i -th cycle under unit transmission times ($T_1 = T_2 = 1$). The process has two states: $(1, 2)$ and $(2, 1)$. Each pair represents an AoI vector, and the number on each arrow represents the selected modality.

In the reformulated SMDP, the policy is a sequence of decisions made at each restart state. Specifically, define a cycle as the period from state $\Delta_{1,\text{re}}$ back to itself. The policy is defined as $\pi_\tau := ((\tau_{1,0}, \tau_{2,0}), (\tau_{1,1}, \tau_{2,1}), \dots)$, where $\tau_{m,i}$ is the decision for state $\Delta_{m,\text{re}}$ in the i -th cycle. Fig. 3 illustrates one cycle of the process in a special case of unit transmission times (i.e., $T_1 = T_2 = 1$). As shown in Fig. 3, the new policy π_τ is derived from the original policy π by grouping transmissions for the same modalities. Although the original policy observes more AoI states than the new policy (such as $(1, 3), \dots, (1, 2 + \tau_{1,i})$), it does not help the scheduling, as the AoI evolution is deterministic given decisions. Hence, the two policies are equivalent, and Problem **OPT** is equivalent to

$$\inf_{\pi_\tau \in \Pi} \limsup_{T \rightarrow \infty} \frac{1}{T} \mathbb{E}_{\pi_\tau} \left[\sum_{t=0}^{T-1} L(\Delta(t)) \right]. \quad (2)$$

According to [19, Proposition 11.4.7.], if the state space is finite and the process always returns to each state within a finite time, then solving Problem (2) is equivalent to solving the Bellman optimality equation. To ensure this, we consider policies that switch to each modality in finite time; that is, there exists a finite number $\tau_{\max} \in \mathbb{N}$ such that $\tau_{m,i} \leq \tau_{\max}$ for all m and i .

We next present the Bellman optimality equation. For each m , the transition probability from $\Delta_{m,\text{re}}$ to $\Delta_{m',\text{re}}$ equals 1. The transition duration is $\tau T_m + T_{m'}$. Let $C_m(\tau)$ denote the transition cost when the decision is τ . For modality 1, we have

$$C_1(\tau) = \sum_{j=1}^{\tau} \sum_{i=0}^{T_1-1} L(T_1 + i, jT_1 + T_2 + i) + \sum_{i=0}^{T_2-1} L(T_1 + i, (\tau + 1)T_1 + T_2 + i),$$

which represents the total inference error during τ transmissions for modality 1 and one transmission for modality 2.

Similarly, for modality 2, we have

$$C_2(\tau) = \sum_{j=1}^{\tau} \sum_{i=0}^{T_2-1} L(T_1 + jT_2 + i, T_2 + i) + \sum_{i=0}^{T_1-1} L(T_1 + (\tau + 1)T_2 + i, T_2 + i).$$

The Bellman optimality equation is given by

$$h(\Delta_{m,\text{re}}) = \min_{\tau \in \{0, 1, \dots, \tau_{\max}\}} [C_m(\tau) - (\tau T_m + T_{m'}) \bar{L}_{\text{opt}} + h(\Delta_{m',\text{re}})], \quad \forall m = 1, 2. \quad (3)$$

While dynamic programming methods (e.g., policy or value iteration) can solve the Bellman optimality equation (3), they suffer from the curse of dimensionality. Therefore, we aim to design a low-complexity policy that is optimal.

B. Optimal Policy

Before presenting the optimal policy, we define the index function. For every $\theta = 0, 1, \dots, \tau_{\max} - 1$, the index function of modality m , denoted by $\gamma_m(\theta)$, is defined as

$$\gamma_m(\theta) := \min_{k \in \{1, 2, \dots, \tau_{\max} - \theta\}} \frac{C_m(\theta + k) - C_m(\theta)}{k T_m}, \quad (4)$$

where the numerator is the additional total inference error and the denominator is the additional transition duration, both when the decision is $\theta + k$ instead of θ (see Remark 1 in Section IV-D for a detailed explanation of the index function).

Theorem 1. Suppose L is uniformly bounded (Assumption 1), $\inf \emptyset = \infty$, and $\theta \in \{0, 1, \dots, \tau_{\max} - 1\}$. For every $\beta \in \mathbb{R}$ and $m \in \{1, 2\}$, define

$$\tau_{m,\text{opt}}(\beta) := \min \{ \min \{ \theta : \gamma_m(\theta) \geq \beta \}, \tau_{\max} \}. \quad (5)$$

The index function γ_m of each modality m is defined in Eq. (4). The optimal stationary policy for Problem **OPT**, represented by (τ_1^*, τ_2^*) , is then given by

$$\tau_m^* = \tau_{m,\text{opt}}(\bar{L}_{\text{opt}}), \quad \forall m \in \{1, 2\}, \quad (6)$$

where the threshold $\bar{L}_{\text{opt}}$ is the optimal objective value of Problem **OPT**. To determine $\bar{L}_{\text{opt}}$, define two functions

$$g_1(\beta) := C_1(\tau_{1,\text{opt}}(\beta)) + C_2(\tau_{2,\text{opt}}(\beta)), \\ g_2(\beta) := (\tau_{1,\text{opt}}(\beta) + 1)T_1 + (\tau_{2,\text{opt}}(\beta) + 1)T_2,$$

and let $g(\beta) := g_1(\beta) - \beta g_2(\beta)$. The threshold $\bar{L}_{\text{opt}}$ is the unique root of the following equation:

$$g(\beta) = 0. \quad (7)$$

We will prove Theorem 1 in Section IV-C. Eq. (5) in Theorem 1 suggests that the optimal stationary policy exhibits an *index-based threshold* structure. That is, the policy selects modality m for θ consecutive transmissions, stopping at the first time when the index function $\gamma_m(\theta)$ exceeds a threshold β ; then, the policy switches to select the other modality. The index function of each modality can be pre-computed

independently, as it depends only on the known parameters (i.e., the inference error function and the transmission times).

Moreover, Eq. (6) shows that, under the optimal policy, both modalities share the same threshold $\bar{L}_{\text{opt}}$, which is exactly the optimal objective value of Problem **OPT**. The threshold $\bar{L}_{\text{opt}}$ can be determined by solving Eq. (7). The following proposition shows that Eq. (7) can be efficiently solved.

Proposition 1. *The function g has the following properties:*

- (i) g is concave, continuous, and strictly decreasing,
- (ii) $\lim_{\beta \rightarrow \infty} g(\beta) = -\infty$ and $\lim_{\beta \rightarrow -\infty} g(\beta) = \infty$.

Proof sketch: It turns out that g is the minimum of multiple linear functions of β with negative coefficients, which leads to the stated properties. The complete proof is similar to that of [1, Lemma 9] and is therefore omitted. ■

Given these properties of g , we can efficiently solve Eq. (7) (i.e., $g(\beta) = 0$) using low-complexity algorithms such as bisection search and Newton's method [20, Algorithm 1-3].

C. Proof of Theorem 1

We prove Theorem 1 by showing that our policy satisfies the Bellman optimality equation (3) and the optimal decision (τ_1^*, τ_2^*) attains the minimum in Eq. (3). According to average-cost SMDP theory [19, Theorem 11.4.4] [21], for an SMDP with countable state and decision space, if a finite scalar $\bar{L}_{\text{opt}}$ and a uniformly bounded function h satisfy the Bellman optimality equation, then an optimal stationary policy exists. Furthermore, a policy is optimal if it attains the minimum in the Bellman optimality equation for all states.

We cannot directly solve each Bellman optimality equation (3) independently, because the optimal objective value $\bar{L}_{\text{opt}}$ appears in the Bellman optimality equation (3) for each state $\Delta_{m,\text{re}}$. To solve this, we fix $\bar{L}_{\text{opt}} = \beta$ for arbitrary $\beta \in \mathbb{R}$ and focus on solving Problem **OPT- β** , defined as

$$\text{OPT-}\beta \quad \min_{\tau \in \{0, 1, \dots, \tau_{\max}\}} [C_m(\tau) - \tau T_m \beta], \quad \forall m \in \{1, 2\}.$$

From the Bellman optimality equation (3) to Problem **OPT- β** , we discard two terms: $T_{m'} \bar{L}_{\text{opt}}$ and $h(\Delta_{m',\text{re}})$, as they are independent of the decision τ . Note that Problem **OPT- β** can be solved separately for each modality. The next proposition shows that the optimal solution for Problem **OPT- β** can be expressed in terms of β .

Proposition 2. *Suppose L is uniformly bounded (Assumption 1), $\inf \emptyset = \infty$, and $\theta \in \{0, 1, \dots, \tau_{\max} - 1\}$. For every $\beta \in \mathbb{R}$ and $m \in \{1, 2\}$, the optimal solution to Problem **OPT- β** is given by $\tau_{m,\text{opt}}(\beta)$ in Eq. (5), i.e.,*

$$\tau_{m,\text{opt}}(\beta) := \min \{ \min \{ \theta : \gamma_m(\theta) \geq \beta \}, \tau_{\max} \}.$$

The index function γ_m of each modality m is defined in Eq. (4).

Proof sketch: Suppose $\tau_{m,\text{opt}}(\beta)$ attains the minimum of Problem **OPT- β** . We use induction to show that for every integer $0 \leq i \leq \tau_{\max}$, if i satisfies Eq. (5), i.e., $\min \{ \min \{ \theta : \gamma_m(\theta) \geq \beta \}, \tau_{\max} \} = i$, then $\tau_{m,\text{opt}}(\beta) = i$. Firstly, we show that $\tau_{m,\text{opt}}(\beta) = 0$ if $\gamma_m(0) \geq \beta$. Assume that for any integer

$1 \leq j \leq \tau_{\max}$, the result holds for $i = 0, 1, \dots, j - 1$, we aim to show the result holds for $i = j$. If $j < \tau_{\max}$, we show that $\tau_{m,\text{opt}}(\beta) = j$ if (i) $\gamma_m(j) < \beta$ for $i = 0, 1, \dots, j - 1$ and (ii) $\gamma_m(j) \geq \beta$; that is, $\tau_{m,\text{opt}}(\beta) = j = \min \{ \theta : \gamma_m(\theta) \geq \beta \}$. If $j = \tau_{\max}$, we show that $\tau_{m,\text{opt}}(\beta) = \tau_{\max} = j$. Combining the two cases yields $\tau_{m,\text{opt}}(\beta) = j$. See our technical report [22] for a detailed proof. ■

Proposition 2 shows that the optimal solution $\tau_{m,\text{opt}}(\beta)$ attains the minimum of the Bellman optimality equation (3) if $\bar{L}_{\text{opt}} = \beta$. Hence, $\tau_{m,\text{opt}}(\bar{L}_{\text{opt}})$ given in Eq. (6) is an optimal policy for Problem **OPT**.

Next, we show that $\bar{L}_{\text{opt}}$ is the unique root of Eq. (7) by constructing a solution to the Bellman optimality equation (3). Given any $\beta \in \mathbb{R}$, we substitute $\bar{L}_{\text{opt}} = \beta$ and $\tau_{m,\text{opt}}(\beta)$ into the Bellman optimality equation (3) and obtain

$$\begin{aligned} h(\Delta_{m,\text{re}}) \\ = C_m(\tau_{m,\text{opt}}(\beta)) - (\tau_{m,\text{opt}}(\beta)T_m + T_{m'})\beta + h(\Delta_{m',\text{re}}), \end{aligned}$$

for each modality m . Let $h(\Delta_{m',\text{re}}) = 0$, we have

$$\begin{aligned} h(\Delta_{m,\text{re}}) &= C_m(\tau_{m,\text{opt}}(\beta)) - (\tau_{m,\text{opt}}(\beta)T_m + T_{m'})\beta, \\ 0 &= C_{m'}(\tau_{m',\text{opt}}(\beta)) - (\tau_{m',\text{opt}}(\beta)T_{m'} + T_m)\beta + h(\Delta_{m,\text{re}}). \end{aligned}$$

By summing the two equations, canceling $h(\Delta_{m,\text{re}})$, and rearranging terms, we obtain an equation of β :

$$\begin{aligned} C_m(\tau_{m,\text{opt}}(\beta)) + C_{m'}(\tau_{m',\text{opt}}(\beta)) \\ - \beta[(\tau_{m,\text{opt}}(\beta) + 1)T_m + (\tau_{m',\text{opt}}(\beta) + 1)T_{m'}] = 0, \end{aligned}$$

which is exactly Eq. (7) in Theorem 1. Finally, we conclude the proof by showing that $\bar{L}_{\text{opt}}$ and $h(\Delta_{m,\text{re}})$ exist and are bounded. Using Proposition 1 and the Intermediate Value Theorem [23], it follows that Eq. (7) has a unique, finite root, i.e., a finite $\bar{L}_{\text{opt}}$ exists. Hence, the optimal decision $\tau_{m,\text{opt}}(\bar{L}_{\text{opt}})$ in Eq. (5) exists, and trivially, its value is bounded by $\tau_{\max}$. Substituting $h(\Delta_{m',\text{re}}) = 0$ and $\tau_m^* = \tau_{m,\text{opt}}(\bar{L}_{\text{opt}})$ into the Bellman optimality equation (3) yields

$$h(\Delta_{m,\text{re}}) = C_m(\tau_m^*) - (\tau_m^*T_m + T_{m'})\bar{L}_{\text{opt}}.$$

Finally, $h(\Delta_{m,\text{re}})$ is bounded because $\bar{L}_{\text{opt}}$, τ_m^* , T_m , and L are all bounded. ■

D. Discussions

We explain the meaning of the index function in Remark 1 and highlight how our index-based threshold policy generalizes and relates to prior work in Remarks 2 and 3.

Remark 1. *The index function $\gamma_m(\theta)$ reflects the **minimum future cost** if the scheduler continues to select modality m after having selected it for θ consecutive transmissions. To see this, consider a symmetric case when $T_1 = T_2 = T_c$ for some constant $T_c \in \mathbb{Z}_+$. Then, the index function of modality 1 in Eq. (5) becomes*

$$\gamma_1(\theta) = \min_k \frac{\sum_{j=1}^k \sum_{i=0}^{T_c-1} L(T_c + i, (\theta + 2 + j)T_c + i)}{kT_c},$$

which calculates the minimum average cost starting from the AoI state $(T_c, (\theta + 3)T_c)$ (i.e., when $j = 1$) until switching.

Furthermore, our index function also handles the asymmetric case. Specifically, suppose $T_1 > T_2$. Define

$$C_a(\theta, k) := \sum_{j=1}^{k-1} \sum_{i=0}^{T_1-1} L(T_1 + i, (\theta + 1 + j)T_1 + T_2 + i),$$

$$C_b(\theta, k) := \sum_{i=0}^{T_2-1} L(T_1 + i, (\theta + 1 + k)T_1 + T_2 + i) \\ + \sum_{i=T_2}^{T_1-1} L(T_1 + i, (\theta + 1)T_1 + T_2 + i).$$

The index function of modality 1 then becomes

$$\gamma_1(\theta) = \min_k \frac{C_a(\theta, k) + C_b(\theta, k)}{kT_1},$$

where $C_a(\theta, k)$ quantifies the total future cost, while $C_b(\theta, k)$ adjusts this cost according to different transmission times.

Remark 2. Due to the non-monotonic AoI functions, our index function is not necessarily monotonic. This result generalizes the two-source scheduling problem in remote estimation [24], where the estimation error is a monotonic function of AoI. Specifically, when the inference error is a non-decreasing function of AoI and $T_1 = T_2 = 1$, the index function of modality 1 reduces to

$$\gamma_1(\theta) = \min_k \frac{\sum_{j=\theta+2}^{\theta+k+1} L(1, j+1)}{k} = L(1, \theta+3),$$

where the last equality holds as the minimum is achieved when $k = 1$ since L is non-decreasing. Then, the optimal policy for modality 1 is $\tau_1^* = \min\{\min\{\theta : L(1, \theta+3) \geq \bar{L}_{\text{opt}}\}, \tau_{\text{max}}\}$, which is equivalent to the result in [24, Proposition 3.5].

Remark 3. A similar index-based threshold policy has been studied in the single-source remote inference setting [1], [13]. Specifically, under stationary random transmission time, the optimal time to send a new feature is when an AoI-based index function exceeds a threshold. Different from the single-source setting, we demonstrate how to design an index-based threshold policy when the modalities are non-separable.

V. NUMERICAL CASE STUDY: ROBOT STATE PREDICTION

In this section, we examine our multimodal remote inference system through a case study on robot state prediction and conduct a trace-driven evaluation to assess our policy.

A. Experimental Setup

Consider a predictor that aims to continuously track the state (e.g., pose and velocity) of a remote robot. The robot gathers both environmental and self-state information using multimodal sensors, such as LiDAR, cameras, and onboard sensors. The predictor infers the robot's state based on features transmitted from the robot. Due to the high dimensionality and varying sizes of features, transmissions often span multiple time slots and differ across modalities.

We use the OpenAI Bipedal Walker as our robot task, where a four-joint robot must run over stumps, pitfalls, and other

obstacles (see [25] for details). The reinforcement learning algorithm used to control the robot is TD3-FORK [26], which achieves state-of-the-art performance on this task. After training the control model, we generate a time-series dataset in the OpenAI Gymnasium simulation environment, consisting of LiDAR rangefinder measurements, robot state information, and joint control signals.

The inference task is to predict the robot's velocities using sequential LiDAR measurements and control signals as two distinct modalities. We adopt the Long Short-Term Memory (LSTM) neural network as the predictor model, due to its widespread use and effectiveness in time-series forecasting. The network architecture includes an input layer, a hidden layer with 20 LSTM cells, and a fully connected output layer. We use 80% of the dataset for training. To incorporate AoI into model training, we augment the dataset as follows: given any AoI vector (δ_1, δ_2) , we construct the feature-label pairs $(X_{1,t-\delta_1}, X_{2,t-\delta_2}; Y_t)$ for all time index t , where the input features are aligned with their corresponding AoI values. The LSTM network takes the AoI vector as two additional input features and is trained on the augmented training dataset.

All experiments were run on a server with an AMD EPYC 7313 CPU (16 cores) and a single NVIDIA A2 GPU.

B. The Impact of AoI on Inference Error

Fig. 4 illustrates how the inference error varies with the AoI for two modalities, with each AoI ranging from 1 to 50. The color intensity or surface height represents the expected inference error on the testing dataset. Although the inference error generally increases with the AoI of either modality, the function is not monotonic. Furthermore, the impact of each modality differs significantly: as the AoI of modality 1 (LiDAR) increases, the inference error grows faster than it does for modality 2 (control signal). It indicates that LiDAR data is more strongly correlated with the target signal. Moreover, the figure suggests that the inference error may not be an additive function of the AoI vector; the effect of one modality's AoI on the inference error depends on the AoI of the other modality. Overall, the results show that the inference error may be a non-monotonic and non-additive function of the AoI vector.

C. Scheduling Policies Evaluation

We compare the following scheduling policies:

- (i) Index-based threshold policy: This is our policy described in Theorem 1. Note that the index function is only needed to determine the optimal solution (τ_1^*, τ_2^*) ; once determined, the scheduler only uses (τ_1^*, τ_2^*) during operation.
- (ii) Round-robin policy: This policy alternates between the two modalities. Notably, this policy is optimal for minimizing the sum of AoI, i.e., $L(\Delta(t)) = \sum_{m=1}^2 \Delta_m(t)$.
- (iii) Uniform random policy: The policy randomly selects one of the two modalities with equal probability.

We use the results obtained in Section V-B as the empirical expected inference error function. We vary the transmission time for each modality, taking values 2, 4, 6, 8, 10, to reflect different feature sizes. Fig. 5 presents the performance of

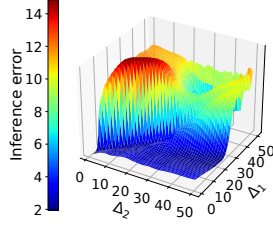


Fig. 4: Inference error vs. AoI

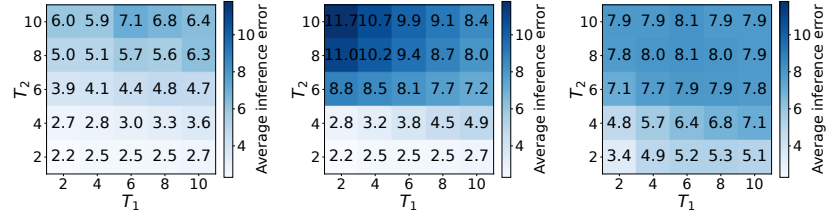


Fig. 5: Performance comparison under varying transmission times

each scheduling policy under varying transmission times. Our proposed index policy consistently achieves the best performance across all cases. Particularly, when $T_1 = 2$ and $T_2 = 6$, our policy reduces the inference error by 55% compared with the round-robin policy and by 53% compared with the uniform random policy. Note that while the round-robin policy is effective for minimizing AoI [15], it fails to capture the complex relationship between inference error and AoI, leading to suboptimal performance in certain scenarios (e.g., when $T_1 = 2$ and $T_2 = 10$). The uniform random policy exhibits similar limitations. These results underscore the importance of jointly considering both modalities and their impact on inference error when designing scheduling policies.

VI. CONCLUSION

We studied the two-modality scheduling problem for remote inference systems, where a receiver-side ML model relies on time-sensitive data from remote sensors. We developed an optimal policy applicable to general AoI penalty functions and heterogeneous transmission times. Experiments on robot state prediction show that our policy reduces inference error by up to 55% compared with baselines. One future extension is to study scheduling over unreliable channels.

REFERENCES

- [1] M. K. C. Shisher, Y. Sun, and I.-H. Hou, "Timely communications for remote inference," *IEEE/ACM Transactions on Networking*, vol. 32, no. 5, pp. 3824–3839, 2024.
- [2] H. Yuan and G. Li, "A survey of traffic prediction: from spatio-temporal data to intelligent transportation," *Data Science and Engineering*, vol. 6, no. 1, pp. 63–85, 2021.
- [3] A. Asvadi, P. Girão, P. Peixoto, and U. Nunes, "3d object tracking using rgb and lidar data," in *IEEE ITSC*, 2016, pp. 1255–1260.
- [4] P. P. Liang, A. Zadeh, and L.-P. Morency, "Foundations & trends in multimodal machine learning: Principles, challenges, and open questions," *ACM Comput. Surv.*, vol. 56, no. 10, pp. 1–42, Jun. 2024.
- [5] A. Li, S. Wu, G. C. Lee, and S. Sun, "From freshness to effectiveness: Goal-oriented sampling for remote decision making," *arXiv preprint arXiv:2504.19507*, 2025.
- [6] M. K. Chowdhury Shisher, H. Qin, L. Yang, F. Yan, and Y. Sun, "The age of correlated features in supervised learning based forecasting," in *IEEE INFOCOM Workshops*, 2021, pp. 1–8.
- [7] M. K. C. Shisher and Y. Sun, "How does data freshness affect real-time supervised learning?" in *ACM MobiHoc*, 2022, p. 31–40.
- [8] S. Kaul, R. Yates, and M. Gruteser, "Real-time status: How often should one update?" in *IEEE INFOCOM*, 2012, pp. 2731–2735.
- [9] R. D. Yates, Y. Sun, D. R. Brown, S. K. Kaul, E. Modiano, and S. Ulukus, "Age of information: An introduction and survey," *IEEE Journal on Selected Areas in Communications*, vol. 39, no. 5, pp. 1183–1210, 2021.
- [10] Y. Sun and B. Cyr, "Sampling for data freshness optimization: Non-linear age functions," *Journal of Communications and Networks*, vol. 21, no. 3, pp. 204–219, 2019.
- [11] A. Maatouk, S. Kriouile, M. Assaad, and A. Ephremides, "The age of incorrect information: A new performance metric for status updates," *IEEE/ACM Trans. Netw.*, vol. 28, no. 5, p. 2215–2228, 2020.
- [12] A. Li, S. Wu, S. Meng, R. Lu, S. Sun, and Q. Zhang, "Toward goal-oriented semantic communications: New metrics, framework, and open challenges," *IEEE Wireless Communications*, vol. 31, no. 5, pp. 238–245, 2024.
- [13] M. K. C. Shisher, B. Ji, I.-H. Hou, and Y. Sun, "Learning and communications co-design for remote inference systems: Feature length selection and transmission scheduling," *IEEE Journal on Selected Areas in Information Theory*, pp. 524–538, 2023.
- [14] Y. Sun, E. Uysal-Biyikoglu, and S. Kompella, "Age-optimal updates of multiple information flows," in *IEEE INFOCOM Workshops*, 2018, pp. 136–141.
- [15] I. Kadota, E. Uysal-Biyikoglu, R. Singh, and E. Modiano, "Minimizing the age of information in broadcast wireless networks," in *54th Allerton*, 2016, pp. 844–851.
- [16] V. Tripathi and E. Modiano, "Optimizing age of information with correlated sources," *IEEE/ACM Transactions on Networking*, vol. 32, no. 6, pp. 4660–4675, 2024.
- [17] B. Zhou and W. Saad, "On the age of information in internet of things systems with correlated devices," in *IEEE GLOBECOM*, 2020, pp. 1–6.
- [18] Y. Sun, E. Uysal-Biyikoglu, R. D. Yates, C. E. Koksal, and N. B. Shroff, "Update or wait: How to keep your data fresh," *IEEE Transactions on Information Theory*, vol. 63, no. 11, pp. 7492–7508, 2017.
- [19] M. L. Puterman, *Markov decision processes: discrete stochastic dynamic programming*. John Wiley & Sons, 2014.
- [20] T. Z. Ornee and Y. Sun, "Sampling and remote estimation for the ornstein-uhlenbeck process through queues: Age of information and beyond," *IEEE/ACM Transactions on Networking*, vol. 29, no. 5, pp. 1962–1975, 2021.
- [21] S. M. Ross, "Average cost semi-markov decision processes," *Journal of Applied Probability*, vol. 7, no. 3, pp. 649–656, 1970.
- [22] K. Zhang, Y. Sun, and B. Ji, "Multimodal remote inference," *arXiv preprint arXiv:2508.07555*, 2025.
- [23] R. G. Bartle and D. R. Sherbert, *Introduction to real analysis*, 4th ed. Wiley, 2011.
- [24] L. Shi and H. Zhang, "Scheduling two gauss-markov systems: An optimal solution for remote state estimation under bandwidth constraint," *IEEE Trans on Signal Processing*, vol. 60, no. 4, pp. 2038–2042, 2012.
- [25] M. Towers, A. Kwiatkowski, J. Terry, J. U. Balis, G. De Cola, T. Deleu, M. Goulão, A. Kallinteris, M. Krimmel, A. KG *et al.*, "Gymnasium: A standard interface for reinforcement learning environments," *arXiv preprint arXiv:2407.17032*, 2024.
- [26] H. Wei and L. Ying, "Fork: A forward-looking actor for model-free reinforcement learning," in *IEEE CDC*, 2021, pp. 1554–1559.