

Online Prediction for Streaming Tensor Time Series

Zhenting Luan, Haoning Wang, Liping Zhang, Shansuo Liang, Wei Han

Abstract—Real-time prediction plays a vital role in various control systems, such as traffic congestion control and wireless channel resource allocation. In these scenarios, the predictor usually needs to track the evolution of the latent statistical patterns in the modern high-dimensional streaming time series continuously and quickly, which presents new challenges for traditional prediction methods. This paper proposes a novel algorithm based on tensor factorization to predict streaming tensor time series online. The proposed algorithm updates the predictor in a low-complexity online manner to adapt to the time-evolving data. Additionally, an automatically adaptive version of the algorithm is presented to mitigate the negative impact of stale data. Simulation results demonstrate that our proposed methods achieve prediction accuracy similar to that of conventional offline tensor prediction methods, while being much faster than them during long-term online prediction. Therefore, our proposed algorithm provides an effective and efficient solution for the online prediction of streaming tensor time series.

Index Terms—Tensor time series, streaming data, online prediction, tensor factorization.

I. INTRODUCTION

A Time series (TS) is a sequence of discrete data points observed over time, which is generated naturally in various areas such as economics, climate, and traffic. The prediction of TS data is crucial in numerous real-world scenarios, including stock market analysis [1] and traffic flow [2]. Conventional methods for TS prediction usually depend on autoregression (AR) models, e.g., autoregressive integrated moving average (ARIMA) model [3] and vector autoregressive (VAR) model [4], to handle scalar or vector batch data.

However, with the emergence of large-scale TS data generated in modern production activities, classical prediction methods face challenges posed by modern sensing technologies. The modern TS data are usually super large-scale and detected from numerous sensors. For instance, the freeway traffic Performance Measurement System (PeMS) of the California Department of Transportation is a sensor network of more than 26000 individual-lane inductive loops installed throughout California freeways, with more than 35000 detectors that send traffic measures to a road control unit per 30 seconds [5], [6]. Classical methods are not efficient and reliable for handling such large-scale TS. Moreover, modern TS data are often multi-dimensional with different characteristics shown in different dimensions, such as the channel state information

in time, spatial, and frequency domains detected by periodic sounding reference signals [7].

With the rapid development of sensing technologies, *streaming TS* data are collected continuously from various types of sensors, such as the traffic flow measurement in the traffic congestion control system [8] and the periodic channel state information (CSI) in wireless channel resource allocation [9]. Accurate real-time prediction is beneficial for the system to make timely decisions in such scenarios. However, conventional prediction methods must re-exploit the entire dataset, including the newly observed data, to rebuild the prediction models in each process of predicting the upcoming data, which is unsuitable for high-frequency streaming TS. Therefore, online prediction is essential to help forecast upcoming data and make immediate decisions continuously and rapidly.

Recent advances in tensor research have led to the development of prediction models for tensor time series (TTS) data with high dimensions based on tensor factorization, such as Tucker decomposition and CP decomposition [10], [11]. These methods explore the joint low-rank structure of all observed data, where the statistical patterns or the periodicity of data are utilized for generating predictors. The tensor factorization strategy is a generalization of matrix factorization, which has been shown to be very efficient in large-scale multi-variable TS [12]. Jing et al. [11] developed a multi-linear prediction model, MOAR, based on Tucker decomposition, which projected the original tensors into several subspaces and built temporal connections among the core tensors by the traditional autoregressive model. The constrained version of MOAR, called MCAR, regularizes the residual of the joint tensor decomposition to find a more stable solution. Shi et al. [13] proposed a prediction algorithm that combined tensor decomposition and ARIMA model for the block-Hankelized TTS data, which exploited the low-rank structure of TTS and captured the intrinsic correlations among it. Both of these approaches demonstrate strong performance for single-time predictions by constructing predictors from scratch with numerous iterative iterations. However, they are not suitable for direct application in streaming TTS online prediction due to the necessity of repeating the process of constructing predictors. Tan et al. [14] presented a short-term traffic flow prediction approach for predicting streaming traffic data based on dynamic tensor completion (DTC). This approach regarded the data to be predicted as some missing data in the corner of a multi-dimensional tensor. It utilized DTC to complete the tensor to obtain the missing data for prediction. However, this process depends on the multi-mode periodicity of traffic flow data, hence is unsuitable for generic TTS data without periodicity.

In this paper, we design an algorithm, TOPA, for predicting generic streaming TTS data in an online manner, which allows

This work was supported in part by the National Natural Science Foundation of China under Grant 12171271. (Corresponding author: Liping Zhang)

Z. Luan, H. Wang, and L. Zhang are with the Department of Mathematical Sciences, Tsinghua University, Beijing 100084, China (e-mail: luanzt18@mails.tsinghua.edu.cn; whn18@mails.tsinghua.edu.cn; lipingzhang@mail.tsinghua.edu.cn)

S. Liang, W. Han are with the Theory Lab, Central Research Institute, 2012 Labs, Huawei Technologies Co., Ltd, Hong Kong, China (e-mail: {liang.shansuo, harvey.hanwei}@huawei.com)

us to quickly update the predictors from the previous model and the latest observations rather than rebuild it from scratch. At each sampling time during the streaming observation, we employ the Tucker decomposition for dimension reduction and joint feature subspace extraction of the streaming TTS, and find the auto-regression pattern in the dimension-reduced TTS (the core tensor series). Then, the prediction results are obtained from the predicted core tensor together with the inverse Tucker decomposition. We design a regularized optimization problem with a least-squares (LS) form to search for the proper Tucker decomposition and regression model. In order to quickly solve the optimization and update the previous predictor, we introduce an online updating scheme for the optimization, of which the initialization is exactly the solution to the previous optimization problem formulated at the last time point. Moreover, for streaming data in which the statistical pattern evolves over a long period, we propose an automatically-adaptive-weight (AAW) strategy for processing past data during online prediction, which can reduce the cumulative error of the predictor from stale data and automatically decrease the weights of very noisy data. In summary, the main contribution of this paper is to propose an online prediction algorithm for streaming TTS based on tensor factorization, which includes the following specific aspects.

- The gradually varied joint tensor factorization structure of streaming TTS is updated by the proposed online updating strategy for tracking the changing underlying dynamic statistical characteristics. Compared with conventional joint tensor factorization methods, this strategy takes much less time owing to only a few alternative iterations with an inherited initialization.
- We design an online updating algorithm with inherited solutions to quickly update the joint tensor factorization structure once observing new data. The prediction accuracy of this algorithm is close to offline methods [11], [13], while the time complexity is significantly lower. The performance of our proposed algorithm is validated in numerical simulations, including synthetic and real-world datasets. We also analyze the convergence of our proposed algorithms.
- To reduce the interference from stale data to the latest statistical characteristics modeling, we propose an automatically adaptive regularization strategy for the introduced LS optimization. This regularization strategy can gradually decrease the weights of historical and abnormal data.

The paper is organized as follows. In Section II, we define the online prediction problem and introduce the joint tensor factorization strategy for TTS. In Section III, we propose the tensor-factorization-based online prediction algorithm and introduce the online manner for predicting streaming TTS data. In order to address the problem of data staleness, we add automatically adaptive weights for regularizing the joint tensor factorization. We analyze the convergence of the proposed algorithms. In Section V, we evaluate the performance of our proposed algorithms in different scenarios and compare it with some tensor offline prediction methods and neural-network-

based methods. Finally, conclusions are drawn in Section VI.

Notations and tensor operators: We employ bold italic lower-case letters, bold italic capital letters, and bold calligraphic letters to denote vectors, matrices, and tensors, respectively (e.g., $\mathbf{x}$, $\mathbf{X}$ and $\mathcal{X}$). For a positive integer n , denote the set $[n] = \{1, 2, \dots, n\}$. We use uppercase letters with script typestyle to denote the set of data (e.g., $\mathcal{G}$). Furthermore, we represent the time series with given length T by uppercase letters with script typestyle equipped with a subscript $[T]$ (e.g., $\mathcal{G}_{[T]}$). Denote $\mathbf{I}_k$ as the $k \times k$ identity matrix. We use $(\cdot)^H$ to denote the conjugate transpose of a complex matrix. The unitary matrix set with given scale $I \times R$ ($I > R$) is denoted as $\mathbb{U}^{I \times R} = \{\mathbf{U} \in \mathbb{C}^{I \times R} | \mathbf{U}^H \mathbf{U} = \mathbf{I}_R\}$. The Frobenius norm of a tensor $\mathcal{X}$ is defined as $\|\mathcal{X}\|_F = \sqrt{\langle \mathcal{X}, \mathcal{X} \rangle}$, where $\langle \cdot, \cdot \rangle$ is the inner product. Moreover, the Frobenius norm of a tuple of tensor $\mathcal{X} = (\mathcal{X}_1, \dots, \mathcal{X}_n)$ is defined as $\|\mathcal{X}\|_F = \sqrt{\sum_{i=1}^n \|\mathcal{X}_i\|_F^2}$. Each dimension of a tensor is called a *mode*. The *m-mode product* $\times_m$ between a tensor $\mathcal{X} \in \mathbb{C}^{I_1 \times \dots \times I_M}$ and a matrix $\mathbf{U}_m \in \mathbb{C}^{J_m \times I_m}$ is denoted as

$$(\mathcal{X} \times_m \mathbf{U}_m)_{i_1 \dots i_{m-1} j_{i_m+1} \dots i_M} = \sum_{i_m=1}^{I_m} \mathcal{X}_{i_1 \dots i_m \dots i_M} \mathbf{U}_{j_{i_m} i_m}, \quad (1)$$

for $\forall j \in [J_m], \forall i_l \in [I_l], \forall l \in [M] \setminus \{m\}$. The *m-mode unfolding matrix* of a tensor is recorded as $\mathcal{X}_{(k)} \in \mathbb{C}^{I_m \times (\prod_{l \neq m} I_l)}$. Denote $R(\mathcal{X}) = (rank(\mathcal{X}_{(1)}), \dots, rank(\mathcal{X}_{(M)}))$ as the *Tucker-rank* of an M th-order tensor $\mathcal{X}$.

II. PRELIMINARIES

A. Problem Formulation

Denote $\mathcal{X}_t \in \mathbb{C}^{I_1 \times \dots \times I_M}$ as the observation at sampling time t and $\mathcal{X}_{[T]} := \{\mathcal{X}_1, \dots, \mathcal{X}_T\}$ as the TTS until sampling time T . Given $\mathcal{X}_{[T]}$, the prediction of $\mathcal{X}_{T+1}$ is to find a predictor $h_T(\cdot)$ that minimizes the mean squared error (MSE) as

$$\min \|\mathcal{X}_{T+1} - h_T(\mathcal{X}_{[T]})\|_F, \quad (2)$$

where $\hat{\mathcal{X}}_{T+1} := h_T(\mathcal{X}_{[T]})$ is the prediction value of $\mathcal{X}_{T+1}$. In this paper, we consider the online prediction problem (2) for streaming TTS, in which $h_T(\cdot)$ changes as streaming observations $\mathcal{X}_T$ arrive sequentially. Therefore, the previous predictor should be updated with the latest observation for the next prediction at each sampling time. This dynamic process is referred to as *online prediction*.

For high-dimensional streaming TTS, the curse of dimensionality makes data processing challenging. To mitigate this challenge and explore the multi-aspect temporal continuity of TTS, we introduce the *joint tensor factorization* strategy in the next subsection, which motivates us to develop low-complexity algorithms for online prediction.

B. Joint Tensor Factorization Strategy

Tensor factorization is a powerful approach for compressing high-dimensional data. In this subsection, we provide a brief overview of this approach, which has been shown to be effective in predicting non-streaming TTS [11].

Owing to the temporal continuity of TTS, the TTS data share similar multi-aspect feature subspaces, which can be

extracted by Tucker decomposition [10]. The Tucker decomposition factorizes a higher-order tensor into a core tensor multiplied by a set of projection matrices along each mode. The core tensor captures the essential information of the original tensor, while the projection matrices represent the feature subspaces. In this paper, we utilize joint Tucker decomposition to compress the original TTS into a lower-dimensional representation while preserving the important features. Specifically, the joint Tucker decomposition finds the $(R_1, \dots, R_M)$ -rank approximation of given TTS with joint projection matrices, which is formulated as:

$$\mathcal{X}_t \approx \mathcal{G}_t^{(T)} \prod_{m=1}^M \times_m \mathbf{U}_m^{(T)}, t \in [T]. \quad (3)$$

where $\mathcal{G}_t^{(T)} \in \mathbb{C}^{R_1 \times \dots \times R_M}$ is the *core tensor* corresponding to $\mathcal{X}_t$, and $\mathbf{U}_m^{(T)} \in \mathbb{U}^{I_m \times R_m}$ is considered as the m -th joint feature subspace of $\mathcal{X}_{[T]}$.

As discussed in [11], the core tensors obtained via joint Tucker decomposition can be viewed as a compact representation of the intrinsic interactions among multiple feature subspaces of TTS. This representation is more suitable for modeling temporal continuity than the original data. Moreover, the size of core tensors is much smaller than the original TTS, which makes it computationally more efficient to apply autoregression models to the core tensor series. For instance, we can utilize AR(p) model to illustrate the temporal correlations among the core tensor series, which is formulated as

$$\mathcal{G}_t^{(T)} = \sum_{i=1}^p \alpha_i^{(T)} \mathcal{G}_{t-i}^{(T)} + \mathcal{E}_t^{(T)}, \quad p+1 \leq t \leq T, \quad (4)$$

where $\{\alpha_i^{(T)}\}$ are AR parameters and $\mathcal{E}_t^{(T)}$ is the white noise in $\mathcal{X}_t$. With (4), the core tensor $\hat{\mathcal{G}}_{T+1}$ of next observation data is predicted by

$$\hat{\mathcal{G}}_{T+1} = \sum_{i=1}^p \alpha_i^{(T)} \mathcal{G}_{T+1-i}^{(T)}, \quad (5)$$

and the prediction of upcoming data is given by extending the temporal continuity of multi-aspect feature subspace to the next time point:

$$\hat{\mathcal{X}}_{T+1} = \hat{\mathcal{G}}_{T+1} \prod_{m=1}^M \times_m \mathbf{U}_m^{(T)}. \quad (6)$$

When dealing with streaming TTS, traditional *offline prediction* methods such as MCAR [11] and BHT-ARIMA [13] suffer from high computational complexity, as they require finding a new predictor for (2) after each observation without inheriting any information from previous predictors. This results in a time-consuming process of repeatedly solving the model during the streaming observation. To address this issue, we propose an **online prediction** method for (2) with the tensor-factorization strategy. The proposed method can quickly find new predictors by updating the previous one with the latest observation, thereby reducing computational complexity.

III. TOPA: TUCKER-DECOMPOSITION-BASED ONLINE PREDICTION ALGORITHM FOR STREAMING TTS

In this section, we introduce a novel two-stage online prediction algorithm for (2) called **Tucker-Decomposition-based Online Prediction Algorithm (TOPA)**. The first stage aims to find an initial predictor using starting TTS. This stage can be seen as a generalization of MCAR [11] with a simplified optimization problem. The second stage is the main focus of TOPA, which continuously updates the predictor and predicts the upcoming data as the streaming observation arrives. We online update the predictor by solving an incremental optimization problem with the previous predictor as initialization. At the end of this section, we analyze the convergence and computational complexity of TOPA.

A. Stage I: Initial Predictor with Starting TTS

In order to find a suitable initial predictor before receiving streaming data, the observer needs to observe a starting TTS as the initial input for TOPA. It is assumed that this starting TTS comprises T_0 samples, and the streaming observation starts from time $T_0 + 1$. By treating the starting TTS as a non-streaming TTS, we utilize the joint tensor-factorization strategy introduced in Section II-B to create the initial predictor.

The joint Tucker decomposition of starting TTS is formulated as (3) with $T = T_0$, which estimates the joint feature subspaces of TTS as $\mathcal{U}^{(T_0)} \triangleq \{\mathbf{U}_m^{(T_0)}\}_{m \in [M]}$ and the core tensor series as $\mathcal{G}_{[T_0]}^{(T_0)} \triangleq \{\mathcal{G}_t^{(T_0)}\}_{t \in [T_0]}$. To be more generalized, we denote $f_{\mathcal{P}^{(T_0)}}(\cdot)$ as any desired AR-type model for $\mathcal{G}_{[T_0]}^{(T_0)}$, e.g., VAR model and ARIMA model, which is formulated as

$$\mathcal{G}_t^{(T_0)} \approx f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{[t-1]}^{(T_0)}), \quad t \in [T_0], \quad (7)$$

where $\mathcal{P}^{(T_0)}$ represents all regression parameters in f and $\mathcal{G}_{[t]}^{(T_0)}$ are the first t core tensors of the starting TTS. For example, $\mathcal{P}^{(T)} = \{\alpha_i^{(T)}\}_i$ in (4).

Remark 1. Notice that the order of the autoregressive model, denoted $p(< T_0)$, defines the number of past data that have direct temporal correlations with the current data, i.e., $f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{[t-1]}^{(T_0)}) = f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{t-p}^{(T_0)}, \dots, \mathcal{G}_{t-1}^{(T_0)})$. For convenience, we still use $f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{[t]}^{(T_0)})$ to represent the model with a little abuse of notation when there is no ambiguity about the order.

With the tensor factorization and regression, TOPA predicts $\mathcal{X}_{T_0+1}$ by inversely projecting the predicted core tensor with the joint feature subspaces:

$$\begin{aligned} \hat{\mathcal{X}}_{T_0+1} &= \hat{\mathcal{G}}_{T_0+1} \prod_{m=1}^M \times_m \mathbf{U}_m^{(T_0)} \\ &= f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{[T_0]}^{(T_0)}) \prod_{m=1}^M \times_m \mathbf{U}_m^{(T_0)} =: h_{T_0}(\mathcal{X}_{[T_0]}), \end{aligned} \quad (8)$$

where $h_{T_0}(\cdot)$ depends on $\mathcal{U}^{(T_0)}$, $\mathcal{G}_{[T_0]}^{(T_0)}$ and $\mathcal{P}^{(T_0)}$. With (8), we can reformulate (2) as the LS optimization problem (9), where the first term in $F_{T_0}(\cdot)$ is used to minimize the error of

core tensor series regression, and the second term is used to regularize the joint Tucker decomposition of the starting TTS with a residual minimization formulation and regularization

$$\begin{aligned} \min_{\mathcal{G}_{[T_0]}^{(T_0)}, \mathcal{P}^{(T_0)}, \mathcal{U}^{(T_0)}} F_{T_0} \left(\mathcal{G}_{[T_0]}^{(T_0)}, \mathcal{P}^{(T_0)}, \mathcal{U}^{(T_0)} \right) &= \sum_{t=p+1}^{T_0} \|\mathcal{G}_t^{(T_0)} - f_{\mathcal{P}^{(T_0)}}(\mathcal{G}_{[t-1]}^{(T_0)})\|_F^2 + \varphi \sum_{t=1}^{T_0} \|\mathcal{G}_t^{(T_0)} - \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_m^{(T_0)})^H\|_F^2 \\ \text{s.t. } \mathbf{U}_m^{(T_0)} &\in \mathbb{U}^{I_m \times R_M}, \forall m \in [M] \end{aligned} \quad (9)$$

According to the multi-variable and LS structure of (9), we propose an iterative algorithm for solving (9) using an alternating minimization approach, as shown in Algorithm 1. In this algorithm, we solve a series of subproblems of (9) in an alternating manner, with each sub-problem fixing all variables except the one being updated. Specifically, in each iteration, we first update the regression model of the core tensor series in terms of current iteration results of core tensors (line 2 in Algorithm 1), and then update the joint projection matrices and core tensor series in the joint Tucker decomposition structure (line 3–6 and line 7–12, respectively). To ensure the convergence of the algorithm, we add a proximal term [15] into each subproblem. The detailed derivation of Algorithm 1 is presented in Appendix A, and the convergence is illustrated at the end of this section.

By executing Algorithm 1, we can build an initial predictor (8) for the starting TTS, which is the basis for the subsequent online prediction stage.

B. Stage II: Online Predictor Updating and Prediction

In this stage, new streaming TTS data arrive in sequence. After each data sampling, we update the previous predictor online and then predict the upcoming TTS data with the updated predictor. The online manner relies on the same assumption of temporal continuity of considered TTS as discussed in Section II-B. In this stage, we repeat *online updating* and *prediction* for online prediction, as illustrated in Fig. 1.

Online updating step: At time $T+1 (> T_0)$, new TTS data $\mathcal{X}_{T+1}$ arrives. Similar to Stage I, the process of finding new predictor $h_{T+1}(\cdot)$ can be formulated as:

$$\begin{aligned} \min_{\mathcal{G}_{[T+1]}^{(T+1)}, \mathcal{P}^{(T+1)}, \mathcal{U}^{(T+1)}} F_{T+1} \left(\mathcal{G}_{[T+1]}^{(T+1)}, \mathcal{P}^{(T+1)}, \mathcal{U}^{(T+1)} \right) &= \quad (10) \\ &\sum_{t=p+1}^{T+1} \|\mathcal{G}_t^{(T+1)} - f_{\mathcal{P}^{(T+1)}}(\mathcal{G}_{[t-1]}^{(T+1)})\|_F^2 \\ &+ \varphi \sum_{t=1}^{T+1} \|\mathcal{G}_t^{(T+1)} - \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_m^{(T+1)})^H\|_F^2 \\ \text{s.t. } \mathbf{U}_m^{(T+1)} &\in \mathbb{U}^{I_m \times R_M}, \forall m \in [M] \end{aligned}$$

Under the assumption of temporal continuity of TTS, the joint feature subspaces among TTS data and core tensor series change smoothly as observation continues [11]. Moreover, compared to $F_T(\cdot)$, $F_{T+1}(\cdot)$ contains only two additional terms that relate to $\mathcal{X}_{T+1}$. Thus the change in the objective function for the online optimization problem is relatively small. When $h_T(\cdot)$ is accurate enough, we can employ it as

parameter $\varphi > 0$.

Algorithm 1 TOPA - Stage I

Input: TTS $\mathcal{X}_{[T_0]}$, and autoregressive model f , step size $\lambda > 0$, tolerance $\epsilon > 0$ and iteration number $k = 0$

Output: Core tensor series $\mathcal{G}_{[T_0]}^{(T_0)}$, autoregressive parameters $\mathcal{P}^{(T_0)}$ and joint projection matrices $\mathcal{U}^{(T_0)}$

Initialize joint projection matrices $\mathcal{U}_0^{(T_0)} = \{\mathbf{U}_{m,0}^{(T_0)}\}_{m=1}^M$ randomly.
Initialize core tensor series:
 $\mathcal{G}_{[T_0],0}^{(T_0)} = \{\mathcal{G}_{t,0}^{(T_0)} := \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,0}^{(T_0)})^H\}_{t=1}^{T_0}$.

- 1: **while** $k \geq 0$ **do**
- 2: Compute regression parameters $\mathcal{P}_{k+1}^{(T_0)}$ for f with $\mathcal{G}_{[T_0],k}^{(T_0)}$ via solving subproblem (21).
- 3: **for** $m = 1 : M$ **do**
- 4: Compute the left and right singular matrices of $\sum_{t=1}^{T_0} \left(\mathcal{X}_t \prod_{i < m} \times_i (\mathbf{U}_{i,k+1}^{(T_0)})^H \prod_{j > m} \times_j (\mathbf{U}_{j,k}^{(T_0)})^H \right)_{(m)} \times \left(\mathcal{G}_{t,k}^{(T_0)} \right)_{(m)}^H + \frac{\lambda}{2\varphi} \mathbf{U}_{m,k}^{(T_0)}$, which denoted as $\mathbf{L}_{m,k+1}$ and $\mathbf{R}_{m,k+1}$, respectively.
- 5: Update $\mathbf{U}_{m,k+1}^{(T_0)} \leftarrow \mathbf{L}_{m,k+1} \mathbf{R}_{m,k+1}^H$
- 6: **end for**
- 7: **for** $t = 1 : p$ **do**
- 8: $\mathcal{G}_{t,k+1}^{(T_0)} \leftarrow \frac{\varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T_0)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T_0)}}{\varphi + \frac{\lambda}{2}}$.
- 9: **end for**
- 10: **for** $t = p+1 : T_0$ **do**
- 11: $\mathcal{G}_{t,k+1}^{(T_0)} \leftarrow \frac{f_{\mathcal{P}_{k+1}^{(T_0)}}(\mathcal{G}_{[t-1],k+1}^{(T_0)}) + \varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T_0)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T_0)}}{1 + \varphi + \frac{\lambda}{2}}$.
- 12: **end for**
- 13: **if** $\|\mathcal{G}_{[T_0],k+1}^{(T_0)} - \mathcal{G}_{[T_0],k}^{(T_0)}\|_F^2 + \|\mathcal{U}_{k+1}^{(T_0)} - \mathcal{U}_k^{(T_0)}\|_F^2 + \|\mathcal{P}_{k+1}^{(T_0)} - \mathcal{P}_k^{(T_0)}\|_F^2 < \epsilon$ **then**
- 14: Break;
- 15: **end if**
- 16: $k \leftarrow k + 1$
- 17: **end while**
- 18: **return** $\mathcal{P}^{(T_0)} = \mathcal{P}_k^{(T_0)}$, $\mathcal{U}^{(T_0)} = \mathcal{U}_k^{(T_0)}$, $\mathcal{G}_{[T_0]}^{(T_0)} = \mathcal{G}_{[T_0],k}^{(T_0)}$

the initialization for solving (10). Specifically, we take $\mathcal{G}_{[T]}^{(T)}$, $\mathcal{U}^{(T)}$ and $\mathcal{P}^{(T)}$ as the initialization of (10), and use an alternative updating scheme to solve (10) in closed forms. The details of alternative updating are shown in the online updating step of Algorithm 2. The derivation of this algorithm is similar to Algorithm 1.

In the practice of our proposed online manner, we only need to run a few iterations to find a predictor accurate enough, owing to the inheritance of the previous solution. In contrast,

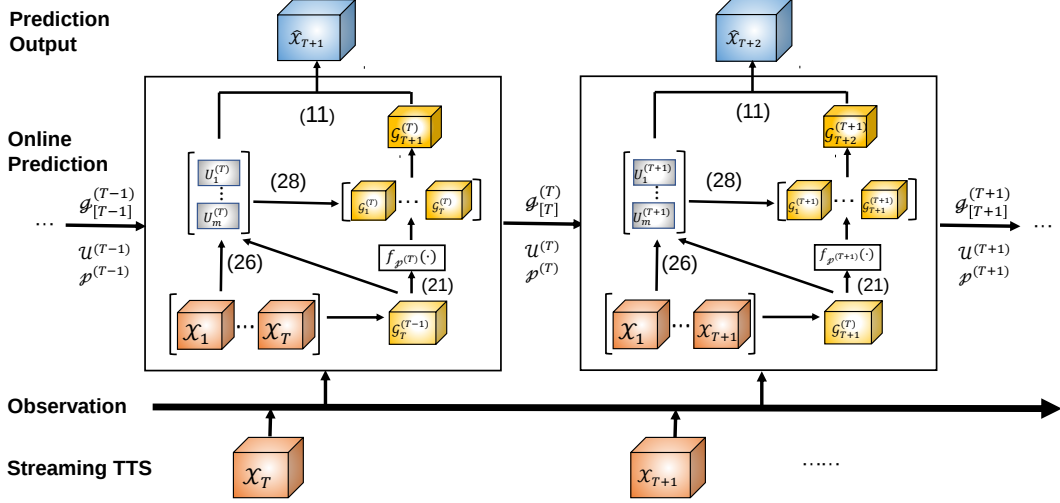


Fig. 1: Stage II of TOPA for Streaming TTS

offline prediction algorithms, such as those discussed in [11] and [13], need to perform multiple iterations with random initialization to solve (10), which is inefficient and repetitive during streaming observations. The simulations in Section V demonstrate that our online updating scheme performs as well as those offline methods, even with just one iteration.

Online prediction step:

After online updating the predictor with new observation, TOPA predicts the upcoming tensor data at time $T + 2$ by

$$\begin{aligned} \hat{\mathcal{X}}_{T+2} &= f_{\mathcal{P}^{(T+1)}} \left(\mathcal{G}_{[T+1]}^{(T+1)} \right) \prod_{m=1}^M \times_m \mathbf{U}_m^{(T+1)} \\ &=: h_{T+1}(\mathcal{X}_{[T+1]}). \end{aligned} \quad (11)$$

C. Convergence and Complexity Analysis

In this subsection, we analyze the convergence and complexity of TOPA. Since the online updating process of Algorithm 2 is similar to Algorithm 1, we only provide the convergence of Algorithm 1 here. The proof of Theorem 1 is presented in Appendix B.

Theorem 1. Assume that the starting TTS and the regression parameter sequence $\{\mathcal{P}_k^{(T_0)}\}_k$ generated by Algorithm 1 are bounded. Then, for any limit points $\mathcal{W}_*^{(T_0)} = (\mathcal{G}_{[T_0],*}^{(T_0)}, \mathcal{P}_*^{(T_0)}, \mathcal{U}_*^{(T_0)})$ of the sequence $\{\mathcal{W}_k^{(T_0)} = (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_k^{(T_0)}, \mathcal{U}_k^{(T_0)})\}_k$, we have

$$\frac{\partial F_{T_0}(\mathcal{W}_*^{(T_0)})}{\partial \mathcal{P}^{(T_0)}} = \mathbf{0}, \quad \frac{\partial F_{T_0}(\mathcal{W}_*^{(T_0)})}{\partial \mathcal{G}_{[T_0]}^{(T_0)}} = \mathbf{0}, \quad (12)$$

and

$$-\frac{\partial F_{T_0}(\mathcal{W}_*^{(T_0)})}{\partial \mathbf{U}_{m,k}^{(T_0)}} \in N(\mathbf{U}_{m,*}^{(T_0)}), \quad (13)$$

where $N(\mathbf{U}_{m,*}^{(T_0)})$ is the normal cone [16] of $\mathbb{U}_{m \times R_m}$ at $\mathbf{U}_{m,*}^{(T_0)}$. In other words, any limit point of the iteration sequence generated by Algorithm 1 is the stationary point of (9).

The computational complexity of stage II of TOPA is dominated by updating joint projection matrices and core tensor series. For each iteration in online prediction, the total computational complexity is $O(2MT\bar{I}^M\bar{R} + M(M-1)T\bar{R}^M\bar{I})$, where we denote $\bar{I} = (\prod_{m \in [M]} I_m)^{1/M}$ and $\bar{R} = (\prod_{m \in [M]} R_m)^{1/M}$. The detailed complexity analysis is presented in Appendix C.

IV. TOPA-AAW: TOPA WITH AUTOMATICALLY ADAPTIVE WEIGHTS

In real-world scenarios, when streaming TTS data are observed continuously over a relatively long period, the effectiveness of stale data in revealing the latest statistical patterns in streaming TTS gradually diminishes. This leads to an increasing cumulative prediction error, such as wireless CSI, in which the intrinsic statistical patterns evolve rapidly as observations continue, and the data becomes outdated quickly. To address this issue, we propose an automatically adaptive weight (AAW) regularization method with a time sliding window τ for modifying (10). The AAW gradually decreases the weights of stale data to reduce their impacts on prediction accuracy. To further eliminate the interference caused by heavily noisy data, we introduce an automatic factor in AAW to reduce the weights of heavily noisy data. The optimization problem for online prediction, once $\mathcal{X}_{T+1}$ is observed at time $T + 1$, is modified as:

Algorithm 2 TOPA - Stage II

Input: Streaming TTS $\mathcal{X}_{[T]}$ from $T = T_0$, autoregressive model f , step size $\lambda > 0$, and tolerance $\epsilon > 0$

Output: Online prediction results $\{\hat{\mathcal{X}}_t\}_{t=T_0+2}^\infty$

```

1: while observing new data  $\mathcal{X}_{T+1}$  at time  $T + 1$ : do
2:    $k \leftarrow 0$ 
3:   step 1: online updating
4:    $\mathcal{G}_{T+1}^{(T)} := \frac{f_{\mathcal{P}^{(T)}}(\mathcal{G}_{[T]}^{(T)}) + \varphi \mathcal{X}_{T+1} \prod_{m=1}^M \times_m (\mathbf{U}_m^{(T)})^H}{1 + \varphi}$ 
5:    $\mathcal{P}_k^{(T+1)} = \mathcal{P}^{(T)}, \mathcal{U}_k^{(T+1)} = \mathcal{U}^{(T)}, \mathcal{G}_{[T+1],k}^{(T+1)} = \mathcal{G}_{[T+1]}^{(T)}$ 
6:   while  $k \geq 0$  do
7:     Compute regression parameters  $\mathcal{P}_{k+1}^{(T+1)}$  for  $f$  with  $\mathcal{G}_{[T+1],k}^{(T+1)}$  via solving subproblem (21).
8:     for  $m = 1 : M$  do
9:       Compute the left and right singular matrices of
         
$$\sum_{t=1}^{T+1} \left( \mathcal{X}_t \prod_{i < m} \times_i (\mathbf{U}_{i,k+1}^{(T+1)})^H \prod_{j > m} \times_j (\mathbf{U}_{j,k}^{(T+1)})^H \right)_{(m)} \left( \mathcal{G}_{t,k}^{(T+1)} \right)_{(m)}^H + \frac{\lambda}{2\varphi} \mathbf{U}_{m,k}^{(T+1)}$$

         which denoted as  $\mathbf{L}_{m,k+1}$  and  $\mathbf{R}_{m,k+1}$ , respectively.
10:      Update  $\mathbf{U}_{m,k+1}^{(T+1)} \leftarrow \mathbf{L}_{m,k+1} \mathbf{R}_{m,k+1}^H$ 
11:    end for
12:    for  $t = 1 : p$  do
13:       $\mathcal{G}_{t,k+1}^{(T+1)} \leftarrow \frac{\varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T+1)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T+1)}}{\varphi + \frac{\lambda}{2}}$ 
14:    end for
15:    for  $t = p + 1 : T + 1$  do
16:       $\mathcal{G}_{t,k+1}^{(T+1)} \leftarrow \frac{f_{\mathcal{P}_{k+1}^{(T+1)}}(\mathcal{G}_{[t-1],k+1}^{(T+1)}) + \varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T+1)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T+1)}}{1 + \varphi + \frac{\lambda}{2}}$ 
17:    end for
18:    if  $\|\mathcal{G}_{[T+1],k+1}^{(T+1)} - \mathcal{G}_{[T+1],k}^{(T+1)}\|_F^2 + \|\mathcal{U}_{k+1}^{(T+1)} - \mathcal{U}_k^{(T+1)}\|_F^2 + \|\mathcal{P}_{k+1}^{(T+1)} - \mathcal{P}_k^{(T+1)}\|_F^2 < \epsilon$  then
19:      Break;
20:    end if
21:     $k \leftarrow k + 1$ 
22:  end while
23:   $\mathcal{P}^{(T+1)} = \mathcal{P}_k^{(T+1)}, \mathcal{U}^{(T+1)} = \mathcal{U}_k^{(T+1)}, \mathcal{G}_{[T+1]}^{(T+1)} = \mathcal{G}_{[T+1],k}^{(T+1)}$ 


---


step 2: online prediction
24:   $\hat{\mathcal{X}}_{T+2} = f_{\mathcal{P}^{(T+1)}}(\mathcal{G}_{[T+1]}^{(T+1)}) \prod_{m=1}^M \times_m \mathbf{U}_m^{(T+1)}$ 
25:  return  $\hat{\mathcal{X}}_{T+2}$ 
26:   $T \leftarrow T + 1$ 
27: end while

```

$$\begin{aligned}
& \min_{\substack{\mathcal{P}^{(T+1)}, \mathcal{U}^{(T+1)}, \\ \mathcal{G}_{[T+1]}^{(T+1)}}} F_{T+1}^{(\text{AAW})}(\mathcal{G}_{[T+1]}^{(T+1)}, \mathcal{P}^{(T+1)}, \mathcal{U}^{(T+1)}) = \quad (14) \\
& \sum_{t=T-\tau+2}^{T+1} \left(\|\mathcal{G}_t^{(T+1)} - f_{\mathcal{P}^{(T+1)}}(\mathcal{G}_{[t-1]}^{(T+1)})\|_F^2 + \right. \\
& \left. \varphi \omega_t^{(T+1)} \|\mathcal{G}_t^{(T+1)} - \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_m^{(T+1)})^H\|_F^2 \right) \\
& \text{s.t. } \mathbf{U}_m^{(T+1)} \in \mathbb{U}^{I_m \times R_m}, \forall m \in [M]
\end{aligned}$$

where $\omega_t^{(T+1)}$ is the AAW

$$\omega_t^{(T+1)} = \begin{cases} (1 - \alpha^{t-(T-\tau+1)}) \max\{\beta, 1 - \epsilon_t^{(T+1)}\}, & T - \tau + 2 \leq t \leq T, \\ 1, & t = T + 1, \end{cases} \quad (15)$$

with Tucker-decomposition residual error

$$\epsilon_t^{(T+1)} = \frac{\|\mathcal{X}_t - \mathcal{G}_t^{(T+1)} \prod_{m=1}^M \times_m \mathbf{U}_m^{(T+1)}\|_F^2}{\|\mathcal{X}_t\|_F^2}. \quad (16)$$

Here $\alpha \in (0, 1)$ is the damping parameter for reducing the accumulative prediction error from stale data. The ‘max’ factor is inversely proportional to the residual error $\epsilon_t^{(T+1)}$ of joint tensor decomposition for streaming TTS in the last prediction, while $\beta \in (0, 1)$ is the minimum residual factor.

Remark 2. The parameter β is configured as a threshold for checking the deviation of the joint Tucker decomposition for each data in TTS, which can help prevent the data with significant decomposition residual errors from being entirely discarded. For stable TTS with gently evolving joint subspaces, such as two real-world datasets in subsection V-B, β can be broadly chosen from 0 to a positive number close to 1 since the threshold is inactive. For unstable TTS with fast-evolving joint subspaces, such as the wireless channel discussed in V-A2, β should be set properly to deal with fast-fading or noisy channel states. In addition, we specifically discuss the effects of choice of α and φ in subsection V-B.

Since the structure of (14) is similar to (10) except for the sliding time window and AAW, we solve (14) in a similar online manner as Algorithm 2. With AAW, the online updating step in Algorithm 2 is modified as follows, while the prediction step is the same.

- **Estimate New Core Tensor** (line 4 in Algorithm 2):

$$\mathcal{G}_{T+1}^{(T)} = \frac{1}{1 + \varphi \omega_{T+1}^{(T+1)}} \left(f_{\mathcal{P}^{(T)}}(\mathcal{G}_{[T]}^{(T)}) + \varphi \omega_{T+1}^{(T+1)} \mathcal{X}_{T+1} \prod_{m=1}^M \times_m (\mathbf{U}_m^{(T)})^H \right). \quad (17)$$

- **Update Projection Matrices** (line 8–11 in Algorithm 2): In $(k + 1)$ -th iteration, for $m = 1$ to M ,

$$\mathbf{U}_{m,k+1}^{(T+1)} = \mathbf{L}_{m,k+1} (\mathbf{R}_{m,k+1})^H,$$

where $\mathbf{L}_{m,k+1}$ and $\mathbf{R}_{m,k+1}$ are the left and right singular matrices of

$$\begin{aligned}
& \sum_{t=T-\tau+2}^{T+1} \omega_t^{(T+1)} \left(\mathcal{X}_t \prod_{i < m} \times_i (\mathbf{U}_{i,k+1}^{(T+1)})^H \prod_{j > m} \times_j (\mathbf{U}_{j,k}^{(T+1)})^H \right)_{(m)} \\
& \times \left(\mathcal{G}_{t,k}^{(T+1)} \right)_{(m)}^H + \frac{\lambda}{2\varphi} \mathbf{U}_{m,k}^{(T+1)}, \quad (18)
\end{aligned}$$

respectively.

- **Update core tensor series** (line 12–17 in Algorithm 2):

In $(k+1)$ -th iteration, for $t = T - \tau + 2$ to $T + 1$,

$$\mathcal{G}_{t,k+1}^{(T+1)} = \frac{1}{1 + \varphi\omega_t^{(T+1)} + \frac{\lambda}{2}} \left[f_{\mathcal{P}_{k+1}^{(T+1)}}(\mathcal{G}_{[t-1],k+1}^{(T+1)}) + \varphi\omega_t^{(T+1)} \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T+1)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T+1)} \right]. \quad (19)$$

With the time sliding window, TOPA-AAW further reduces the computational complexity to τ/T times that of TOPA, as Appendix C analyzes.

V. NUMERICAL EXPERIMENTS

In this section, we evaluate the performance of TOPA in various scenarios, including synthetic low-rank TTS, wireless channel simulations, and two real-world datasets. All numerical experiments are conducted using MATLAB R2018b on a Windows PC with a quad-core Intel(R) Core(TM) 2.0GHz CPU and 8 GB RAM.

To the best of our knowledge, TOPA is the first method for online prediction of generic streaming TTS. In order to provide a reference for comparing the performance of TOPA, we also conduct five other non-online prediction methods:

- **BHT-ARIMA [13] and MCAR [11]:** These are two effective offline one-shot prediction methods for TTS. We run these offline methods using the same online observations at each sampling time as the online methods by building their predictors from scratch with multiple iterations to forecast the next tensor data. BHT-ARIMA employs BHT, a Hankelized tensor expansion technique, to further exploit the temporal relationships among TTS and uses a structure similar to MCAR to formulate predictors. The offline methods perform sufficient iterations until convergence in each prediction, thereby revealing good prediction performance.
- **TOPA-init:** To illustrate the importance of the online updating for predictors, we test the performance of TOPA without model updating in some experiments. In other words, we use the initial predictor obtained from Algorithm 1 and the streaming observations to forecast upcoming data at each sampling time.
- **LSTM [17] and GMRL [18]:** These are two neural-network-based prediction methods. Long Short-Term Memory (LSTM) is a classical time-series forecasting method in the deep learning area, which utilizes a type of recurrent neural network (RNN) architecture that is designed to process and retain information over long sequences. GMRL is a latest work based on neural networks equipped with the tensor decomposition framework.

When comparing the time costs of LSTM and GMRL, we did not account for the training time costs of neural networks, though it may cost much time in practice. On the other hand, owing to the predictor inheritance of TOPA, we conduct TOPA and TOPA-AAW with only one iteration at each sampling time during the online prediction.

In addition, taking two real-world datasets in subsection V-B as examples, we discuss the effects of parameter choice for TOPA-AAW. We evaluate the prediction accuracy of

algorithms with the Normalized Root Mean Square Error (NRMSE) metric.

A. Synthetic Datasets

1) *Synthetic Low-Rank Streaming TTS:* We generate a low-rank noisy streaming TTS with a similar process as [19, Sec. 6.3]. The low-rank TTS structure follows the joint Tucker decomposition. We first generate the core tensor series $\{\mathcal{G}_t \in \mathbb{C}^{4 \times 4 \times 4}\}_{t=1}^{70}$ with ARIMA(3, 1, 0) model. Then, we generate the joint feature subspace matrices $\{\mathbf{U}_m \in \mathbb{U}^{20 \times 4}\}_{m=1}^3$, by orthogonalizing randomly generated matrices with i.i.d. $\mathcal{N}(0, 1)$ entries. The generated noisy TTS are formulated as

$$\mathcal{X}_t = \mathcal{G}_t \prod_{m=1}^M \times_m \mathbf{U}_m + \rho \|\mathcal{G}_t\|_F \mathcal{E}_t \in \mathbb{C}^{20 \times 20 \times 20}, \quad (20)$$

where $\rho = 0.1$ is the noise parameter and $\mathcal{E}_t$ is the noise tensor with i.i.d. $\mathcal{N}(0, 1)$ entries. Note that we have $\|\mathcal{G}_t\|_F = \|\mathcal{G}_t \prod_{m=1}^M \times_m \mathbf{U}_m\|_F$ and ρ represents the ratio of noise in TTS.

The generated TTS with noise is divided into two sets: the first $T_0 = 20$ data in the training set, and the rest 50 in the test set. The regularization parameters for TOPA-AAW are set as $\varphi = 10, \alpha = 0.4, \beta = 0.4$. Fig. 2 compares the average NRMSE of different prediction methods with 10000 Monte Carlo experiments, each independently and randomly generates a TTS as (20). The ‘‘TOPA-init’’ line in Fig. 2 shows that prediction error increases as prediction continues without model online updating due to the noises in TTS. With the online updating manner, TOPA reveals almost the same performance as MCAR, while BHT-ARIMA has a little advantage over these two methods. Two neural-network-based methods do not perform as well as other algorithms, due to the implicit regression patterns of the core tensor series. Among all these methods, TOPA and TOPA-AAW are much more efficient owing to the one-loop algorithms, while also keep high prediction accuracy similar to the offline methods.

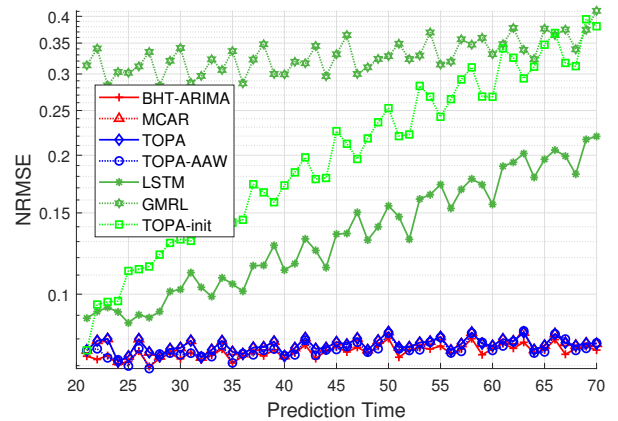


Fig. 2: Performance of different prediction methods on synthetic TTS (20). TOPA and TOPA-AAW reveal almost the same performance as the offline methods.

TABLE I: Average time costs and NRMSE of different algorithms on synthetic TTS

	TOPA	TOPA-AAW	MCAR	BHT-ARIMA	GMRL	LSTM
Time(ms)	47.8	25.1	77.4	238.4	200.7	72.2
NRMSE	0.0776	0.0760	0.0778	0.0750	0.3299	0.1390

2) *Wireless CSI Prediction*: We consider the wireless channel prediction problem with high-frequency observation and wild fluctuation of statistical features. For fifth-generation wireless communication, many massive multiple-input multiple-output (MIMO) transmission techniques highly rely on accurate CSI acquisition, such as precoding and beamforming. However, the observation of CSI and later data processing usually quickly become outdated due to the short coherent time. CSI prediction is a natural way to tackle this issue. For time-varying channels, the online manner is important to predict CSI in upcoming time slots.

We use QuaDRiGa [20] to generate realistic CSI streaming data for single base-station (BS) and single user-equipment (UE) downlink transmission. BS and UE are equipped with 16 and 2 antennas, respectively. The transmission occupies the 700MHz frequency with 10MHz bandwidth and 50 subcarriers. During the CSI observation, BS is static, and UE moves with 1.5m/s speed. The CSI observation frequency is configured as $f_s = 25\text{Hz}$, so we sample CSI per 40ms. Then, the CSI observation produces a streaming TTS $\{\mathcal{X}_t \in \mathbb{C}^{16 \times 2 \times 50}\}_{t=1}^{\infty}$, of which $(\mathcal{X}_t)_{ijk}$ denotes the channel state at time t on k -th subcarrier between BS's i -th antenna and UE's j -th antenna.

We run 1000 times Monte Carlo (MC) simulations to evaluate the performance of our proposed algorithms. For each MC simulation, we randomly select a continuous part of CSI TTS with 70 samples and let the first $T_0 = 20$ CSI tensors be the training set, while the last 50 tensors are the streaming data observed in the following 50 sampling time. The scale of core tensors is set as $(10, 2, 20)$, which significantly compresses the TTS data. The regularization parameters for the residual of tensor decomposition are set as $\varphi = 10, \alpha = 0.9, \beta = 0.6$. Since the CSI fluctuates wildly, we shorten the time sliding window in TOPA-AAW to $\tau = 8$. All prediction methods employ ARIMA(2, 1, 1) model for building temporal relations among core tensor series. To further improve the adaptability of TOPA/TOPA-AAW for the time-varying property of CSI, we conduct two iterations instead in the online updating step of Algorithm 2.

TABLE II: Average Time costs and NRMSE of different algorithms on CSI prediction

	TOPA	TOPA-AAW	MCAR	BHT-ARIMA	GMRL	LSTM
Time(ms)	186.8	37.6	915.8	2634.7	22.6	50.7
NRMSE	0.0636	0.0591	0.0635	0.0634	0.1295	0.9170

Fig. 3 illustrates the average NRMSE of CSI online prediction for different prediction methods. The “TOPA-init” line in Fig. 3 reveals the extreme change in the wireless environment, even in low-speed mobility scenarios. As an online method, TOPA shows very close performance to two offline methods, while the prediction accuracy of TOPA-AAW is much better than TOPA and MCAR, owing to the good adaptation of AAW

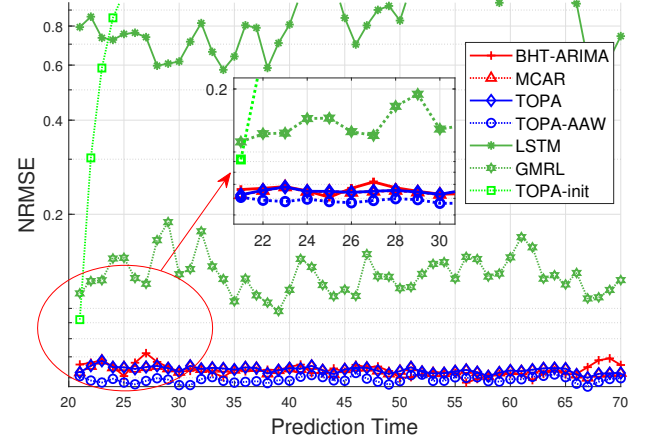


Fig. 3: Performance of different prediction methods on CSI prediction. TOPA-AAW has the best performance, while TOPA performs as well as the offline methods.

to the wireless channel evolution.

Table II presents the average time costs and NRMSE of six methods in each prediction. It should be emphasised that the training of the GMRL and LSTM methods cost up to a few minutes, which is not considered in Table II. TOPA-AAW is much more efficient than other methods except GMRL, while has the best prediction performance. Table II illustrates that the difference in the number of necessary alternative iterations for building predictors leads to a significant difference in the time costs between online and offline methods.

B. Real-world Datasets: USHCN and NASDAQ100

In this subsection, We apply our prediction algorithm to two real-world datasets:

- *USHCN*¹: this dataset records the monthly climate data of 1218 weather stations in the United States during the past 100 years, including four statistical features: monthly mean maximum temperature, monthly mean minimum temperature, monthly minimum temperature, and monthly total precipitation. 120 meteorological stations with relatively complete data are screened for numerical experiments. With their quarterly average observation data from 1940 to 2014, we establish a 75-length TTS with the scale $120 \times 4 \times 4$ of each tensor data.
- *NASDAQ100*²: This dataset records the daily business data of 102 NASDAQ-listed companies over 90 days in 2014, including five statistical features: opening quotation, highest quotation, lowest quotation, closing quotation and adjusting the closing price. Therefore, we obtain a 90-length TTS with the scale 102×5 of each matrix data.

We divide the screened USHCN TTS into two parts: the first $T_0 = 40$ data are used for finding the initial predictor and predicting the data at sampling time 41, and the last 35 data

¹<https://www.ncei.noaa.gov/pub/data/ushcn/v2.5/>

²<https://github.com/Karin-Karlsson/stockdata>

are regarded as the streaming observation of USHCN TTS. In terms of the discussion in [11, Sec. V], we choose AR(2) model to build the temporal correlations among core tensor series. In the structure of joint tensor decomposition, the scale of core tensors is given as $12 \times 4 \times 4$. Since the statistical characteristics of data are nearly stable, we configure a wide time sliding window in TOPA-AAW with length $\tau = 20$.

For the NASDAQ100 dataset, we find the initial predictor with the first $T_0 = 60$ traffic flow data, and streamly observe the last 30 data. The ARIMA(3, 1, 0) model captures the statistical characteristics in NASDAQ100 TTS. The time sliding window used in TOPA-AAW is configured as $\tau = 20$.

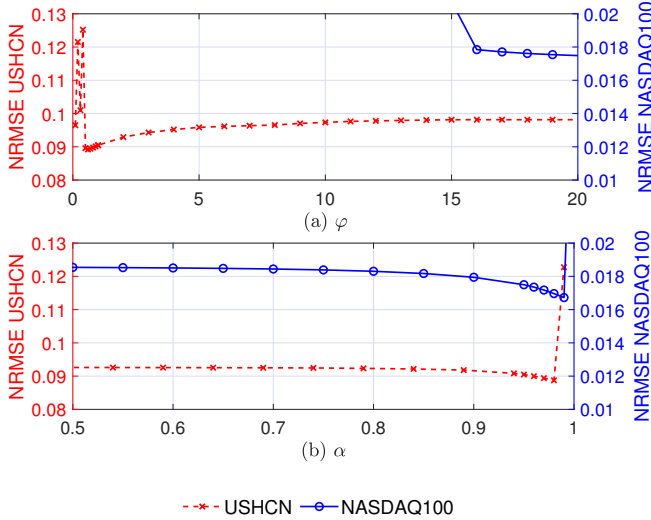


Fig. 4: Effect of parameter φ and α for the TOPA-AAW algorithm on two real-world datasets: (a) average NRMSE with $\alpha = 0.95$ and different φ , and (b) average NRMSE with $\varphi_{\text{USHCN}} = 0.5$, $\varphi_{\text{NASDAQ100}} = 20$ and different α . When $\alpha < 1$ and φ are large enough, the accuracy of TOPA-AAW is not sensitive to the choice of these parameters.

We first research the effects of different choices of α and φ in (14). Fig. 4 shows the average NRMSE of the proposed TOPA-AAW algorithm with different values of parameters φ and α . Regarding Fig. 4 (a), we can see that the best φ of USHCN and NASDAQ100 datasets are 0.2 and 20, respectively. For both two datasets, the prediction performance is unsatisfactory when φ is relatively small, which manifests the necessity of regularization. For stable TTS, small φ makes the regularization term close to 0 with sufficiently accurate joint Tucker decomposition, hence the objective function in (14) is dominated by the first term, which leads to inaccurate joint decomposition with online prediction continuing. Furthermore, when φ is large enough, the accuracy of TOPA-AAW is improved to a stable level. From Fig. 4 (b), the best α for USHCN and NASDAQ100 datasets is close to 1, which illustrates the effectiveness of reducing the weight of stale data. The prediction performance is extremely poor when α is set as 1, since the regularization in (14) with $\alpha = 1$ only retains the last term with time $t = T + 1$. In brief, when $\alpha < 1$ and φ are

large enough, the accuracy of TOPA-AAW is not very sensitive to the choice of these parameters, which allows us to choose the parameter more flexibly. In the following experiments, we configure the regularization parameters for TOPA-AAW as: $\varphi = 0.5$ and $\alpha = 0.98$ for USHCN, and $\varphi = 20$ and $\alpha = 0.99$ for NASDAQ100. Moreover, as discussed in Remark 2, we set $\beta = 0.5$.

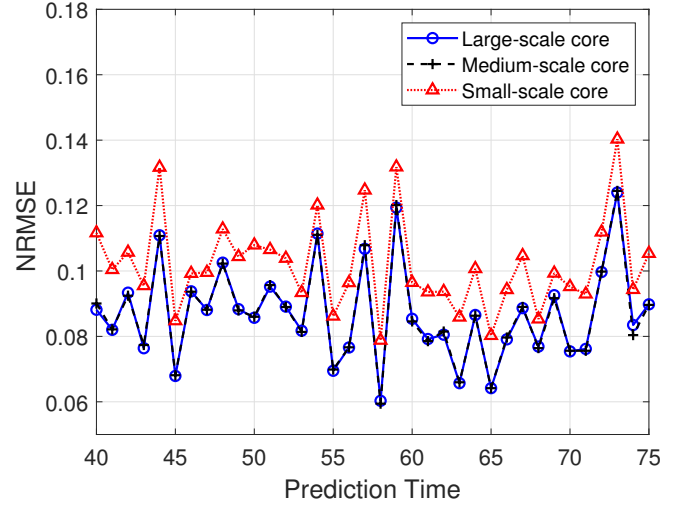


Fig. 5: Performance of different sizes of core tensors for TOPA-AAW algorithm on USHCN dataset.

TABLE III: Average time costs of different sizes of core tensors for TOPA-AAW algorithm on USHCN dataset

	Large scale	Medium scale	Small scale
Time(ms)	8.7	5.1	4.8

Fig. 5 and Table III show the average NRMSE and time costs of the proposed TOPA-AAW algorithm with different scales of core tensors, respectively. We implement the experiments with three different sizes of the core tensors: large scale with dimensions of $80 \times 4 \times 4$, medium scale with dimensions of $12 \times 4 \times 4$, and small scale with dimensions of $3 \times 3 \times 3$. The results show that a larger scale of core tensors leads to higher prediction accuracy with more time costs. Therefore, choosing the right size of core tensor is a trade-off between prediction accuracy and time costs.

Next we focus on the impact of the length of the time sliding window on the prediction performance of TOPA-AAW. As shown in Fig. 6, a longer sliding time window (with $\tau = 5$ and $\tau = 10$) leads to better prediction performance. On the other side, the complexity analysis in Section III-C proposes that the computational complexity of TOPA-AAW is positively proportional to the time window length τ . With the trade-off between prediction accuracy and time costs, we choose the length of time sliding window $\tau = 20$ in the following experiments.

Fig. 7 depicts the comparison results of prediction performance. TOPA reveals nearly the same accuracy for two real-world datasets as offline methods, though it runs much

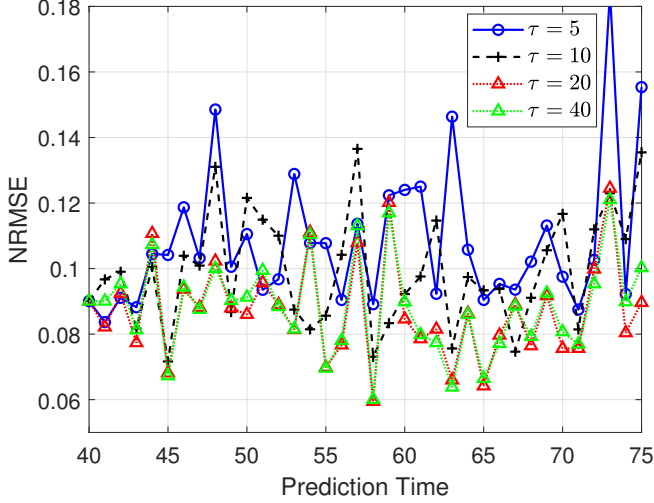


Fig. 6: Performance of different sliding time window lengths for TOPA-AAW algorithm on USHCN dataset.

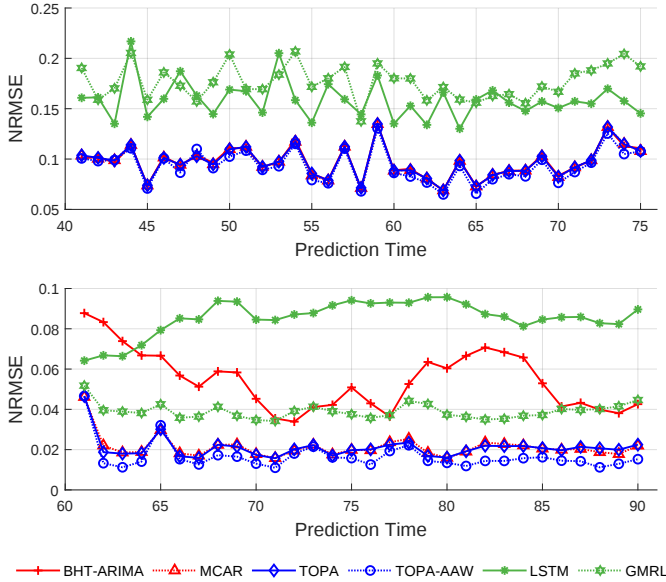


Fig. 7: Performance of different prediction methods on two real-world datasets: (a) USHCN, and (b) NASDAQ100. TOPA-AAW shows the best performance owing to AAW.

TABLE IV: Average time costs of different algorithms on USHCN and NASDAQ100 datasets

Time(ms)	TOPA	TOPA-AAW	MCAR	BHT-ARIMA	GMRL	LSTM
USHCN	59.2	5.1	328.7	1583.1	15.0	34.6
NASDAQ100	53.5	11.4	263.0	419.2	6.9	26.2

fewer iterations than MCAR in each prediction. TOPA-AAW can further improve the accuracy of TOPA by tracking the latest statistical patterns in TTS. In the NASDAQ100 dataset, the prediction accuracy of TOPA-AAW is nearly 20% better than TOPA and MCAR, while three methods show significant advantages over BHT-ARIMA and two neural-network-based methods. Table IV presents the average time costs of six pre-

diction methods. Owing to the online manner, TOPA/TOPA-AAW can provide prediction results efficiently and accurately.

VI. CONCLUSION

In this paper, we present a joint-tensor-factorization-based online prediction algorithm for streaming tensor time series. By leveraging tensor factorization, our algorithm effectively compresses the streaming data while capturing the underlying intrinsic correlations. The proposed online updating scheme significantly enhances the speed of predictor updating and can maintain high prediction accuracy. We also analyze the convergence of the proposed algorithm. Additionally, we introduce automatically adaptive weights to address the challenge of data staleness in streaming data. Through the numerical experiments in various scenarios, we observe promising results that validate the effectiveness and efficiency of our approach.

APPENDIX A

DERIVATION OF THE ALTERNATIVE ITERATIONS IN TOPA

The derivation of Algorithm 1 relies on the alternative manner and proximal term added in each subproblem. In each iteration, Algorithm 1 alternatively updates only one of $\mathcal{P}^{(T_0)}$, $\mathcal{U}^{(T_0)}$, and $\mathcal{G}_{[T_0]}^{(T_0)}$ at a time. In each subproblem during alternative updating, we add a proximal term to ensure the convergence of the algorithm. Since the process of Algorithm 2 is similar to Algorithm 1, we only need to present the derivation of Algorithm 1. For the sake of convenience, we take the AR model (4) for f as an example to illustrate the derivation process.

- **Update regression parameters:** For AR model, the subproblem of updating $\mathcal{P}^{(T_0)} = \{\alpha^{(T_0)} = [\alpha_1^{(T_0)}, \dots, \alpha_p^{(T_0)}]\}$ with proximal term and step size λ is formulated as

$$\begin{aligned} \alpha_{k+1}^{(T_0)} &= \arg \min_{\alpha} F_{T_0} \left(\mathcal{G}_{[T_0],k}^{(T_0)}, \{\alpha\}, \mathcal{U}_k^{(T_0)} \right) + \frac{\lambda}{2} \|\alpha - \alpha_k^{(T_0)}\|_F^2 \\ &= \arg \min_{\alpha} \sum_{t=p+1}^{T_0} \|\mathcal{G}_{t,k}^{(T_0)} - \sum_{i=1}^p \alpha_i \mathcal{G}_{t-i,k}^{(T_0)}\|_F^2 + \frac{\lambda}{2} \|\alpha - \alpha_k^{(T_0)}\|_F^2. \end{aligned} \quad (21)$$

The closed-form solution to (21) is

$$\alpha_{k+1}^{(T_0)} = \left(\mathbf{R} + \frac{\lambda}{2} \mathbf{I}_p \right)^{-1} \left(\mathbf{r} + \frac{\lambda}{2} \alpha_k^{(T_0)} \right), \quad (22)$$

where

$$\mathbf{R}_{i,j} = \sum_{t=p+1}^{T_0} \left\langle \mathcal{G}_{t-i,k}^{(T_0)}, \mathcal{G}_{t-j,k}^{(T_0)} \right\rangle, 1 \leq i, j \leq p, \quad (23)$$

and

$$\mathbf{r}_i = \sum_{t=p+1}^{T_0} \left\langle \mathcal{G}_{t-i,k}^{(T_0)}, \mathcal{G}_{t,k}^{(T_0)} \right\rangle, 1 \leq i \leq p, \quad (24)$$

Remark 3. When T_0 is large enough, the entries of $\mathbf{R}$ and $\mathbf{r}$ can be seemed as approximations of the autocorrelation function of $\mathcal{G}_{[T_0],k}^{(T_0)}$:

$$r(j-i) = \mathbb{E} \langle \mathcal{G}_t, \mathcal{G}_{t+j-i} \rangle.$$

Then, we have

$$\mathbf{R}_{i,j} \approx r(i-j), \quad \mathbf{r}_i \approx r(i).$$

With these approximations, when $\lambda = 0$, (22) is the solution to the known Yule-Walker equation [21].

- **Update joint projection matrices** $\mathcal{U}^{(T+1)}$: For $m = 1$ to M , the subproblem of updating $\mathbf{U}_m^{(T_0)}$ with proximal term and step size λ is formulated as

$$\begin{aligned} & \mathbf{U}_{m,k+1}^{(T_0)} \\ &= \arg \min_{\mathbf{U} \in \mathbb{U}^{I_m \times R_m}} F_{T_0} \left(\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_{k+1}^{(T_0)}, \right. \\ & \quad \left. \{\mathbf{U}_{i,k+1}^{(T_0)}\}_{i=1}^{m-1} \cup \{\mathbf{U}\} \cup \{\mathbf{U}_{j,k}^{(T_0)}\}_{j=m+1}^M \right) \\ & \quad + \frac{\lambda}{2} \|\mathbf{U} - \mathbf{U}_{m,k}^{(T_0)}\|_F^2, \\ &= \arg \min_{\mathbf{U} \in \mathbb{U}^{I_m \times R_m}} \varphi \sum_{t=1}^{T_0} \|\mathcal{G}_{t,k}^{(T_0)} - \mathcal{H}_t \times_m \mathbf{U}^H\|_F^2 + \frac{\lambda}{2} \|\mathbf{U} - \mathbf{U}_{m,k}^{(T_0)}\|_F^2 \\ & \stackrel{a}{=} \arg \max_{\mathbf{U} \in \mathbb{U}^{I_m \times R_m}} \Re \text{trace} \left(\mathbf{U}^H \left(\sum_{t=1}^{T_0} (\mathcal{H}_t)_{(m)} (\mathcal{G}_{t,k}^{(T_0)})_{(m)}^H + \frac{\lambda}{2\varphi} \mathbf{U}_{m,k}^{(T_0)} \right) \right) \end{aligned} \quad (25)$$

where $\mathcal{H}_t = \mathcal{X}_t \prod_{i < m} \times_i (\mathbf{U}_{i,k+1}^{(T_0)})^H \prod_{j > m} \times_j (\mathbf{U}_{j,k}^{(T_0)})^H$, and $\stackrel{a}{=}$ utilizes the orthogonal constraints $\mathbf{U}^H \mathbf{U} = \mathbf{I}_{R_m}$.

With orthogonal constraints, we have

$$\mathbf{U}_{m,k+1}^{(T_0)} = \mathbf{L}_{m,k+1} \mathbf{R}_{m,k+1}^H, \quad (26)$$

where $\mathbf{L}_{m,k+1} \in \mathbb{U}^{I_m \times R_m}$ and $\mathbf{R}_{m,k+1} \in \mathbb{U}^{R_m \times R_m}$ are the left and right singular matrices of

$$\sum_{t=1}^{T_0} (\mathcal{H}_t)_{(m)} (\mathcal{G}_{t,k}^{(T_0)})_{(m)}^H + \frac{\lambda}{2\varphi} \mathbf{U}_{m,k}^{(T_0)}. \quad (27)$$

As m traverses $[M]$, $\mathcal{U}_k^{(T_0)}$ is updated to $\mathcal{U}_{k+1}^{(T_0)} \triangleq \{\mathbf{U}_{m,k+1}^{(T_0)}\}_{m \in [M]}$.

- **Update core tensor series** $\mathcal{G}_{[T_0]}^{(T_0)}$: For $t = 1$ to T_0 , the subproblem of updating $\mathbf{U}_m^{(T_0)}$ with proximal term and step size λ is formulated as

$$\begin{aligned} \mathcal{G}_{t,k+1}^{(T_0)} &= \arg \min_{\mathcal{G}_t} F_{T_0} (\mathcal{G}_{[t-1],k+1}^{(T_0)} \cup \{\mathcal{G}_t\} \cup \\ & \quad \{\mathcal{G}_{s,k}^{(T_0)}\}_{s=t+1}^{T+1}, \mathcal{P}_{k+1}^{(T_0)}, \mathcal{U}_{k+1}^{(T_0)}) \\ & \quad + \frac{\lambda}{2} \|\mathcal{G}_t - \mathcal{G}_{t,k}^{(T_0)}\|_F^2 \end{aligned} \quad (28)$$

Since (28) is a quadratic optimization problem, when the gradient of the objective function with respect to $\mathcal{G}_t$ vanishes, we can obtain the solution to (28):

$$\mathcal{G}_{t,k+1}^{(T_0)} = \begin{cases} \frac{1}{\varphi + \frac{\lambda}{2}} [\varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T_0)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T_0)}], \\ \frac{1}{1 + \varphi + \frac{\lambda}{2}} [f_{\mathcal{P}_{k+1}^{(T_0)}}(\mathcal{G}_{[t-1],k+1}^{(T_0)}) + \\ \varphi \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k+1}^{(T_0)})^H + \frac{\lambda}{2} \mathcal{G}_{t,k}^{(T_0)}], \end{cases}$$

As t traverses $[T+1]$, $\mathcal{G}_{[T_0],k}^{(T_0)}$ is updated to $\mathcal{G}_{[T_0],k+1}^{(T_0)} := \{\mathcal{G}_{t,k+1}^{(T_0)}\}_{t \in [T]}$.

APPENDIX B PROOF OF THEOREM 1

Denote $\mathcal{W}_k^{(T_0)} = \{\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_k^{(T_0)}, \mathcal{U}_k^{(T_0)}\}$. In terms of the alternative manner of Algorithm 1, we have

$$\begin{aligned} & F_{T_0} (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_k^{(T_0)}, \mathcal{U}_k^{(T_0)}) \\ & \stackrel{b}{\geq} F_{T_0} (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_{k+1}^{(T_0)}, \mathcal{U}_k^{(T_0)}) + \frac{\lambda}{2} \|\mathcal{P}_{k+1}^{(T_0)} - \mathcal{P}_k^{(T_0)}\|_F^2 \\ & \stackrel{c}{\geq} F_{T_0} (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_{k+1}^{(T_0)}, \mathcal{U}_{k+1}^{(T_0)}) + \\ & \quad \frac{\lambda}{2} \|\mathcal{P}_{k+1}^{(T_0)} - \mathcal{P}_k^{(T_0)}\|_F^2 + \frac{\lambda}{2} \|\mathcal{U}_{k+1}^{(T_0)} - \mathcal{U}_k^{(T_0)}\|_F^2 \\ & \stackrel{d}{\geq} F_{T_0} (\mathcal{G}_{[T_0],k+1}^{(T_0)}, \mathcal{P}_{k+1}^{(T_0)}, \mathcal{U}_{k+1}^{(T_0)}) + \frac{\lambda}{2} \|\mathcal{W}_{k+1}^{(T_0)} - \mathcal{W}_k^{(T_0)}\|_F^2, \end{aligned} \quad (29)$$

where $\stackrel{b}{\geq}, \stackrel{c}{\geq}, \stackrel{d}{\geq}$ are due to the optimization subproblem (21), (25), (28), respectively. Therefore, via recursion, we have

$$\begin{aligned} & F_{T_0} (\mathcal{G}_{[T_0],0}^{(T_0)}, \mathcal{P}_0^{(T_0)}, \mathcal{U}_0^{(T_0)}) \\ & \geq F_{T_0} (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_k^{(T_0)}, \mathcal{U}_k^{(T_0)}) + \sum_{i=1}^k \frac{\lambda}{2} \|\mathcal{W}_i^{(T_0)} - \mathcal{W}_{i-1}^{(T_0)}\|_F^2. \end{aligned} \quad (30)$$

For any k and $t \leq T_0$, with the boundedness of $\mathcal{X}_t$ and $\mathbf{U}_{m,k}^{(T_0)}$, we have

$$\begin{aligned} \|\mathcal{G}_{t,k}^{(T_0)}\|_F^2 &\leq \|\mathcal{G}_{t,k}^{(T_0)} - \mathcal{X}_t \prod_{m=1}^M \times_m (\mathbf{U}_{m,k}^{(T_0)})^H\|_F^2 + 2\|\mathcal{X}_t\|_F^2 \\ &\leq 2F_{T_0} (\mathcal{G}_{[T_0],k}^{(T_0)}, \mathcal{P}_k^{(T_0)}, \mathcal{U}_k^{(T_0)}) + 2\|\mathcal{X}_t\|_F^2 \\ &\leq 2F_{T_0} (\mathcal{G}_{[T_0],0}^{(T_0)}, \mathcal{P}_0^{(T_0)}, \mathcal{U}_0^{(T_0)}) + 2\|\mathcal{X}_t\|_F^2. \end{aligned} \quad (31)$$

Hence $\{\mathcal{G}_{[T_0],k}^{(T_0)}\}_k$ is bounded. Therefore, $\{\mathcal{W}_k^{(T_0)}\}$ is bounded and has at least one limit point.

With (30), we have

$$\sum_{i=1}^{+\infty} \|\mathcal{W}_i^{(T_0)} - \mathcal{W}_{i-1}^{(T_0)}\|_F^2 < \infty, \quad (32)$$

which implies

$$\mathcal{W}_{k+1}^{(T_0)} - \mathcal{W}_k^{(T_0)} \rightarrow \mathbf{0}. \quad (33)$$

For any limit point $\mathcal{W}_*^{(T_0)} = (\mathcal{G}_{[T_0],*}^{(T_0)}, \mathcal{P}_*^{(T_0)}, \mathcal{U}_*^{(T_0)})$ of $\{\mathcal{W}_k^{(T_0)}\}$, there exists a subsequence $\{\mathcal{W}_{k_i}^{(T_0)}\}_i$ that converges to $\mathcal{W}_*^{(T_0)}$. With (33), we have $\mathcal{W}_{k_i+1}^{(T_0)} \rightarrow \mathcal{W}_*^{(T_0)}$.

For any t and i , since $\mathcal{G}_{t,k_i+1}^{(T_0)}$ is the solution to (28) with index k_i , we have

$$\begin{aligned} & \frac{\partial F_{T_0}}{\partial \mathcal{G}_t} (\mathcal{G}_{[t-1],k_i+1}^{(T_0)} \cup \{\mathcal{G}_{t,k_i+1}^{(T_0)}\} \cup \\ & \quad \{\mathcal{G}_{s,k_i}^{(T_0)}\}_{s=t+1}^{T+1}, \mathcal{P}_{k_i+1}^{(T_0)}, \mathcal{U}_{k_i+1}^{(T_0)}) + \lambda (\mathcal{G}_{t,k_i+1}^{(T_0)} - \mathcal{G}_{t,k_i}^{(T_0)}) = \mathbf{0}. \end{aligned} \quad (34)$$

Let $i \rightarrow +\infty$, (34) converges to

$$\frac{\partial F_{T_0}(\mathcal{W}_*^{(T_0)})}{\partial \mathcal{G}_t^{(T_0)}} = \mathbf{0}. \quad (35)$$

Similarly, we have

$$\frac{\partial F_{T_0}(\mathcal{W}_*^{(T_0)})}{\partial \mathcal{P}^{(T_0)}} = \mathbf{0}. \quad (36)$$

Then, (12) is proved.

For any m and i , since $\mathbf{U}_{m,k_i+1}^{(T_0)}$ is the solution to (25) with index k_i , with [16, Theorem 8.15] we have

$$-\frac{\partial F_{T_0}}{\partial \mathbf{U}_m^{(T_0)}}(\mathcal{G}_{[T_0],k_i}^{(T_0)}, \mathcal{P}_{k_i+1}^{(T_0)}, \{\mathbf{U}_{i,k_i+1}^{(T_0)}\}_{i=1}^{m-1} \cup \{\mathbf{U}_{m,k_i+1}^{(T_0)}\} \cup \{\mathbf{U}_{j,k_i}^{(T_0)}\}_{j=m+1}^M) \\ - \lambda(\mathbf{U}_{m,k_i+1}^{(T_0)} - \mathbf{U}_{m,k_i}^{(T_0)}) \in N(\mathbf{U}_{m,k_i+1}^{(T_0)}). \quad (37)$$

Let $i \rightarrow +\infty$, with [16, Proposition 6.6], (37) converges to

$$-\frac{\partial F_{T_0}}{\partial \mathbf{U}_m^{(T_0)}}(\mathcal{W}_*^{(T_0)}) \in N(\mathbf{U}_{m,*}^{(T_0)}). \quad (38)$$

Then (13) is proved.

APPENDIX C

COMPLEXITY ANALYSIS FOR TOPA AND TOPA-AAW

We list the computational complexity of each iteration during online prediction by TOPA in Table V. Denote $\bar{I} = (\prod_{m \in [M]} I_m)^{1/M}$ and $\bar{R} = (\prod_{m \in [M]} R_m)^{1/M}$ as the geometric average of the scale of TTS data and core tensors, respectively. Here we denote the total order of autoregressive model f is p , which is assumed much less than the scale of tensor data.

TABLE V: Computational Complexity of Stage II in TOPA

Step	Computational Complexity
Line 4 of TOPA-Stage II	$O(M\bar{I}^M \bar{R})$
Computing $\mathcal{P}_{k+1}^{(T+1)}$	$O(p^3 \bar{R}^M)$
(26)	$O(MT((M-1)\bar{I}^M \bar{R} + \bar{R}^M \bar{I}))$
(28)	$O(T(M\bar{I}^M \bar{R} + p\bar{R}^M))$
Total	$O(M^2 T \bar{I}^M \bar{R} + M T \bar{R}^M \bar{I})$

Owing to the time sliding window, TOPA-AAW reduces the computational complexity to τ/T times that of TOPA, with $O(M^2 \tau \bar{I}^M \bar{R} + M \tau \bar{R}^M \bar{I})$.

REFERENCES

- [1] A. Sharma, D. Bhuriya, and U. Singh, "Survey of stock market prediction using machine learning approach," in *2017 International conference of Electronics, Communication and Aerospace Technology (ICECA)*, vol. 2, 2017, pp. 506–509.
- [2] Y. Lv, Y. Duan, W. Kang, Z. Li, and F.-Y. Wang, "Traffic flow prediction with big data: A deep learning approach," *IEEE Transactions on Intelligent Transportation Systems*, vol. 16, no. 2, pp. 865–873, 2015.
- [3] J. D. Hamilton, *Time series analysis*. Princeton university press, 2020.
- [4] E. Zivot and J. Wang, "Vector autoregressive models for multivariate time series," *Modeling financial time series with S-PLUS*, pp. 385–429, 2006.
- [5] A. Pascale and M. Nicoli, "Adaptive bayesian network for traffic flow prediction," in *2011 IEEE Statistical Signal Processing Workshop (SSP)*, 2011, pp. 177–180.
- [6] X. Chen and L. Sun, "Bayesian temporal factorization for multidimensional time series prediction," *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 44, no. 9, pp. 4659–4673, 2021.
- [7] C. Luo, J. Ji, Q. Wang, X. Chen, and P. Li, "Channel State Information Prediction for 5G Wireless Communications: A Deep Learning Approach," *IEEE Transactions on Network Science and Engineering*, vol. 7, no. 1, pp. 227–236, 2020.
- [8] A. I. Maarala, M. Rautiainen, M. Salmi, S. Pirttikangas, and J. Riekk, "Low latency analytics for streaming traffic data with apache spark," in *2015 IEEE International Conference on Big Data*. IEEE, 2015, pp. 2855–2858.
- [9] K. B. Letaief and Y. J. Zhang, "Dynamic multiuser resource allocation and adaptation for wireless systems," *IEEE Wireless Communications*, vol. 13, no. 4, pp. 38–47, 2006.
- [10] T. G. Kolda and B. W. Bader, "Tensor decompositions and applications," *SIAM Review*, vol. 51, no. 3, pp. 455–500, 2009.
- [11] P. Jing, Y. Su, X. Jin, and C. Zhang, "High-order temporal correlation learning for time-series prediction," *IEEE transactions on cybernetics*, vol. 49, no. 6, pp. 2385–2397, 2018.
- [12] H.-F. Yu, N. Rao, and I. S. Dhillon, "Temporal regularized matrix factorization for high-dimensional time series prediction," *Advances in neural information processing systems*, vol. 29, 2016.
- [13] Q. Shi, J. Yin, J. Cai, A. Cichocki, T. Yokota, L. Chen, M. Yuan, and J. Zeng, "Block hankel tensor ARIMA for multiple short time series forecasting," *AAAI 2020 - 34th AAAI Conference on Artificial Intelligence*, pp. 5758–5766, 2020.
- [14] H. Tan, Y. Wu, B. Shen, P. J. Jin, and B. Ran, "Short-term traffic prediction based on dynamic tensor completion," *IEEE Transactions on Intelligent Transportation Systems*, vol. 17, no. 8, pp. 2123–2133, 2016.
- [15] A. Kaplan and R. Tichatschke, "Proximal point methods and nonconvex optimization," *Journal of global Optimization*, vol. 13, pp. 389–406, 1998.
- [16] R. T. Rockafellar and R. J.-B. Wets, *Variational analysis*. Springer Science & Business Media, 2009, vol. 317.
- [17] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural computation*, vol. 9, no. 8, pp. 1735–1780, 1997.
- [18] J. Deng, J. Deng, R. Jiang, and X. Song, "Learning gaussian mixture representations for tensor time series forecasting," *arXiv preprint arXiv:2306.00390*, 2023.
- [19] Y. Sun, Y. Guo, C. Luo, J. Tropp, and M. Udell, "Low-rank tucker approximation of a tensor from streaming data," *SIAM Journal on Mathematics of Data Science*, vol. 2, no. 4, pp. 1123–1150, 2020.
- [20] S. Jaeckel, L. Raschkowski, K. Börner, and L. Thiele, "QuaDRiGa: A 3-D multi-cell channel model with time evolution for enabling virtual field trials," *IEEE transactions on antennas and propagation*, vol. 62, no. 6, pp. 3242–3256, 2014.
- [21] W. Chen, B. D. Anderson, M. Deistler, and A. Filler, "Solutions of yule-walker equations for singular ar processes," *Journal of Time Series Analysis*, vol. 32, no. 5, pp. 531–538, 2011.