

WORKING PAPER · NO. 2023-47

Prediction When Factors are Weak

Stefano Giglio, Dacheng Xiu, and Dake Zhang

MARCH 2023

Prediction When Factors are Weak*

Stefano Giglio[†]

Yale School of Management

NBER and CEPR

Dacheng Xiu[‡]

Booth School of Business

University of Chicago

Dake Zhang[§]

Booth School of Business

University of Chicago

This Version: March 28, 2023

Abstract

In macroeconomic forecasting, principal component analysis (PCA) has been the most prevalent approach to the recovery of factors, which summarize information in a large set of macro predictors. Nevertheless, the theoretical justification of this approach often relies on a convenient and critical assumption that factors are pervasive. To incorporate information from weaker factors, we propose a new prediction procedure based on supervised PCA, which iterates over selection, PCA, and projection. The selection step finds a subset of predictors most correlated with the prediction target, whereas the projection step permits multiple weak factors of distinct strength. We justify our procedure in an asymptotic scheme where both the sample size and the cross-sectional dimension increase at potentially different rates. Our empirical analysis highlights the role of weak factors in predicting inflation, industrial production growth, and changes in unemployment.

Key words: Supervised PCA, PCA, PLS, weak factors, marginal screening

*We benefited tremendously from discussions with seminar and conference participants at Gregory Chow Seminar Series in Econometrics and Statistics, Glasgow University, University of Science and Technology of China, Nankai University, IAAE Invited Session at ASSA 2023, Conference on Big Data and Machine Learning in Econometrics, Finance, and Statistics at the University of Chicago, and 14th Annual SoFiE Conference.

[†]Address: 165 Whitney Avenue, New Haven, CT 06520, USA. E-mail address: stefano.giglio@yale.edu.

[‡]Address: 5807 S Woodlawn Avenue, Chicago, IL 60637, USA. E-mail address: dacheng.xiu@chicagobooth.edu.

[§]Address: 5807 S Woodlawn Avenue, Chicago, IL 60637, USA. Email: dkzhang@chicagobooth.edu.

1 Introduction

Starting from the seminal contribution of [Stock and Watson \(2002\)](#), factor models have played a prominent role in macroeconomic forecasting. Principal component analysis (PCA), advocated in that paper, has been the most prevalent approach to the recovery of factors that summarize the information contained in a large set of macroeconomic predictors, and reduce the dimensionality of the forecasting problem.

The theoretical justification for the PCA approach to factor analysis often relies on a convenient – but critical – assumption that factors are pervasive (*strong*), see for example [Bai and Ng \(2002\)](#) and [Bai \(2003\)](#). In that case, the common components of predictors can be extracted consistently by PCA and separated from the idiosyncratic components. Recently, [Bai and Ng \(2021\)](#) relax this condition, showing that PCA can consistently recover the underlying factors under weaker assumptions.

Nevertheless, PCA is an unsupervised approach, and by its nature, this poses some limits to its ability to find the most useful low-dimensional predictors in a forecasting context. Specifically, if the signal-to-noise ratio is sufficiently low, the factor space spanned by the principal components is inconsistent, or even nearly orthogonal to the space spanned by true factors, see [Hoyle and Rattray \(2004\)](#) and [Johnstone and Lu \(2009\)](#). In such instances, we refer to the underlying factors as *weak*.

In this paper we study a setting in which factors are sufficiently weak that PCA fails to recover them. We propose a new approach to dimension reduction for forecasting, based on *supervised PCA* (SPCA). The key idea of supervised PCA is to select a subset of predictors that are correlated with the prediction target before applying PCA. The concept of supervised PCA originated from a cancer diagnosis technique applied to DNA microarray data by [Bair and Tibshirani \(2004\)](#), and was later formalized by [Bair et al. \(2006\)](#) in a prediction framework, in which some predictors are not correlated with the latent factors that drive the outcome of interest. [Bai and Ng \(2008\)](#) generalize this selection procedure (i.e., a form of hard-thresholding) to what they call the use of targeted predictors (that include soft-thresholding as well), and find it helpful in a macroeconomic forecasting environment.

Unlike [Bair et al. \(2006\)](#), our supervised PCA proposal involves an additional projection step, and a subsequent iterative procedure over selection, PCA, and projection to extract latent factors. More specifically: we first select a subset of the predictors that correlate with the target, and extract a first factor from that subset using PCA. Then, we project the target and all the predictors (including those not selected) on the first factor, and take the residuals. We then repeat the selection step using these residuals, extract a second factor from the new subset using PCA, and then project again the residuals of the target and all predictors on this second factor. We keep iterating these steps until all factors are extracted, each from a different subset of predictors (or their residuals). We provide examples to illustrate that our iterative procedure is necessary in general settings where factors can grow at distinct rates (that is, they are of different strength) and factors are not necessarily marginally correlated with the target. The final step of our procedure is to make predictions with estimated factors via time-series regressions.

We justify our procedure in an asymptotic scheme where both the sample size and the cross-sectional dimension increase but at potentially different rates. We show that our iterative procedure delivers consistent prediction of the target. While our procedure can extract weak factors, we do not have asymptotic guarantee for recovery of the factor space that is orthogonal to the target. Importantly, this is irrelevant for consistency in prediction. Intuitively, using information about the correlation between each predictor and the target, we gain additional information useful to extract some of the factors even when they are weak. As a result, the factor space that we may fail to recover must be orthogonal to the target, and therefore missing it does not affect the consistency of the prediction.

The weak factor problem in our setting arises from the factor loading matrix, whose eigenvalues increase but at a potentially slower rate than the cross-sectional dimension. The factors we consider are weaker than those discussed in [Bai and Ng \(2021\)](#); as we show in the paper, PCA cannot consistently recover them, and prediction via PCA is biased. Interestingly, in this setting even supervised procedures may in general fail to recover the relevant factors: specifically, we show that a widely used supervised procedure, partial least squares (PLS), is in fact subject to the same bias as PCA. That said, our procedure will miss factors that are *extremely weak*. These are the kind of factors studied by [Onatski \(2009\)](#) and [Onatski \(2010\)](#), cases in which the eigenvalues corresponding to the factor component are of the same order of magnitude as those of the idiosyncratic component. In this context, while it is possible to infer the number of factors, [Onatski \(2012\)](#) show that the factor space cannot be recovered consistently (and neither SPCA will be able to do so).

Finally, beyond consistency (which requires weaker assumptions), if we make an additional assumption that each of the latent factors is correlated with at least one of the variables in a multivariate target, we can obtain stronger results: we can estimate the number of weak factors consistently, recover the space spanned by all factors, as well as provide a valid prediction interval on the target. Our asymptotic result does not rely on a perfect recovery of the set of predictors that are correlated with the factors, unlike [Bair et al. \(2006\)](#). Moreover, our result accounts for potential errors accumulated over the iterative procedure.

We illustrate the use of SPCA with an empirical application to macroeconomic forecasting (in the Online Appendix). We combine the standard Fred-Md dataset of 127 macroeconomic and financial variables with the Blue Chip Financial Forecasts dataset, that contains hundreds of forecasts of various variables (like interest rates and inflation) from professional forecasters, thus obtaining a large dataset of predictors. We then apply different prediction and dimension reduction methods to forecast quarterly inflation, industrial production growth, and changes in unemployment. We compare the results using SPCA to those obtained using PCA (as in [Stock and Watson \(2002\)](#)) and PLS (as in [Kelly and Pruitt \(2013\)](#)). We show that in a setting with a large number of (potentially noisy and/or redundant) predictors, SPCA performs well in forecasting macroeconomic quantities out of sample. We also investigate the selection that SPCA operates, and find that it isolates, for each target, a different group of useful predictors; it also focuses on a few financial forecasters, whose

survey responses are selected particularly often. Finally, we illustrate the use of SPCA with multiple targets at the same time (macroeconomic variables forecasted at different horizons: 1, 2, 3, 6 and 12 months).

Our paper relates to several strands of the literature on forecasting and on dimension reduction. Within the context of forecasting using latent factors, it focuses on static approximate factor models. Dynamic factor models are developed in [Forni et al. \(2000\)](#), [Forni and Lippi \(2001\)](#), [Forni et al. \(2004\)](#), and [Forni et al. \(2009\)](#), in which the lagged values of the unobserved factors may also affect the observed predictors. It is possible to extend our approach to the dynamic factor setting, which is beyond the scope of this paper. [Chao and Swanson \(2022\)](#) study estimation and forecasting within a weak-factor-augmented VAR framework. They also use a pre-selection step since factors only have influence on a subset of predictors. A unique contribution of theirs is a self-normalized score statistics for selection in place of correlation screening as in supervised PCA, which ensures consistent selection of marginally correlated predictors with vanishing Type I and II errors. Similar to [Bair et al. \(2006\)](#), they assume all factors to have the same order of strength and all important predictors to be marginally correlated with the target, which our iterative procedure is designed to avoid.

Our paper is also related to a strand of the literature on spike covariance models defined in [Johnstone \(2001\)](#), where the largest few eigenvalues in the covariance matrix differ from the rest in population, yet are still bounded. In this setting, [Bai and Silverstein \(2006\)](#), [Johnstone and Lu \(2009\)](#) and [Paul \(2007\)](#) show that the largest sample eigenvalues and their corresponding eigenvectors are inconsistent unless the sample size grows at a faster rate than the increase of the cross-sectional dimension. [Wang and Fan \(2017\)](#) extend this setting to the case of diverging eigenvalue spikes, and characterize the limiting distribution of the extreme eigenvalues and certain entries of the eigenvectors in a regime where the sample size grows much slower than the dimension. All these papers shed light on the source of bias with the standard PCA procedure in various asymptotic settings.

Besides supervised PCA, an alternative route taken by an adjacent literature to resolving the inconsistency of PCA is sparse PCA, which imposes sparsity on population eigenvectors, see, e.g., [Jolliffe et al. \(2003\)](#), [Zou et al. \(2006\)](#), [d’Aspremont et al. \(2007\)](#), [Johnstone and Lu \(2009\)](#), and [Amini and Wainwright \(2009\)](#). [Uematsu and Yamagata \(2021\)](#) adopt a variant of the sparse PCA algorithm proposed in [Uematsu et al. \(2019\)](#) to estimate a sparsity-induced weak factor model. Because sparsity is rotation dependent, such weak factor models require rotation-specific identification assumptions, whereas standard factor models do not. The weak factor models we consider, for instance, avoid such a sparsity assumption, which makes our approach distinct from the sparse PCA.

In a companion paper, [Giglio et al. \(2020\)](#) incorporate a similar supervised PCA algorithm into the two-pass cross-sectional regression and study its parameter inference in an asset pricing context. In contrast, this paper focuses on predictive inference and provides asymptotic guarantee on the convergence of extracted factors. Our approach also shares the spirit with [Bai and Ng \(2008\)](#) and [Huang et al. \(2021\)](#). The former suggests a hard or soft thresholding procedure to select “targeted”

predictors to which PCA is then applied, without providing theoretical justification. The latter suggests scaling each predictor with its predictive slope on the prediction target before applying the PCA. Our procedure and its asymptotic justification are more involved because the eigenvalues of the factor loadings in our setting can grow at distinct and slower rates.

The rest of the paper is organized as follows. In Section 2 we introduce the model, provide examples to illustrate the impact of weak factors on prediction, and develop our supervised PCA procedure. In Section 3, we present our approach in general settings and provide asymptotic theory for our procedure. Section 4 provides Monte Carlo simulations demonstrating the finite-sample performance. Section 5 concludes. The appendix provides mathematical proofs of the main theorems in the paper. The online appendix presents an empirical application, details on the asymptotic variance estimation, and proofs of propositions and technical lemmas.

2 Methodology

2.1 Notation

Throughout the paper, we use (A, B) to denote the concatenation (by columns) of two matrices A and B . For any time series of vectors $\{a_t\}_{t=1}^T$, we use the capital letter A to denote the matrix $(a_1, a_2, \dots, a_T)$, $\bar{A}$ for $(a_{1+h}, a_{2+h}, \dots, a_T)$, and $\underline{A}$ for $(a_1, a_2, \dots, a_{T-h})$, for some h . We use $[N]$ to denote the set of integers: $\{1, 2, \dots, N\}$. For an index set $I \subset [N]$, we use $|I|$ to denote its cardinality. We use $A_{[I]}$ to denote a submatrix of A whose rows are indexed in I .

We use $a \vee b$ to denote the max of a and b , and $a \wedge b$ as their min for any scalars a and b . We also use the notation $a \lesssim b$ to denote $a \leq Kb$ for some constant $K > 0$ and $a \lesssim_P b$ to denote $a = O_P(b)$. If $a \lesssim b$ and $b \lesssim a$, we write $a \asymp b$ for short. Similarly, we use $a \asymp_P b$ if $a \lesssim_P b$ and $b \lesssim_P a$.

We use $\lambda_{\min}(A)$ and $\lambda_{\max}(A)$ to denote the minimum and maximum eigenvalues of A , and use $\lambda_i(A)$ to denote the i -th largest eigenvalue of A . Similarly, we use $\sigma_i(A)$ to denote the i th singular value of A . We use $\|A\|$ and $\|A\|_F$ to denote the operator norm (or ℓ_2 norm), and the Frobenius norm of a matrix $A = (a_{ij})$, that is, $\sqrt{\lambda_{\max}(A'A)}$, and $\sqrt{\text{Tr}(A'A)}$, respectively. We also use $\|A\|_{\text{MAX}} = \max_{i,j} |a_{ij}|$ to denote the ℓ_∞ norm of A on the vector space. We use $\mathbb{P}_A = A(A'A)^{-1}A'$ and $\mathbb{M}_A = \mathbb{I}_d - \mathbb{P}_A$, for any matrix A with d rows, where $\mathbb{I}_d$ is a $d \times d$ identity matrix.

2.2 Model Setup

Our objective is to predict a $D \times 1$ vector of targets, y_{T+h} , h -step ahead from a set of N predictor variables x_t with a sample of size T .

We assume that x_t follows a linear factor model, that is,

$$x_t = \beta f_t + \beta_w w_t + u_t, \quad (1)$$

where f_t is a $K \times 1$ vector of latent factors, w_t is an $M \times 1$ vector of observed variables, u_t is an $N \times 1$ vector of idiosyncratic errors satisfying $E(u_t) = 0$, $E(f_t u_t') = 0$, and $E(w_t u_t') = 0$. Without loss of generality, we also impose that $E(f_t w_t') = 0$.¹

We assume that the target variables in y are related to x through factors f in a predictive model:

$$y_{t+h} = \alpha f_t + \alpha_w w_t + z_{t+h}, \quad (2)$$

where z_{t+h} is a $D \times 1$ vector of prediction errors.

Using the aforementioned notation, we can rewrite the above two equations in their matrix form as

$$\begin{aligned} X &= \beta F + \beta_w W + U, \\ \bar{Y} &= \alpha \underline{F} + \alpha_w \underline{W} + \bar{Z}. \end{aligned}$$

We now discuss assumptions that characterize the data generating processes (DGPs) of these variables. For clarity of the presentation, we use high-level assumptions, which can easily be verified by standard primitive conditions for i.i.d. or weakly dependent series. Our asymptotic analysis assumes that $N, T \rightarrow \infty$, whereas h, K, D , and M are fixed constants.

Assumption 1. *The factor F , the prediction error Z , and the observable regressor W , satisfy:*

$$\begin{aligned} \|T^{-1} \underline{F} \underline{F}' - \Sigma_f\| &\lesssim_P T^{-1/2}, \|F\|_{\text{MAX}} \lesssim_P (\log T)^{1/2}, \|T^{-1} \underline{W} \underline{W}' - \Sigma_w\| \lesssim_P T^{-1/2}, \\ \|\underline{W} \underline{F}'\| &\lesssim_P T^{1/2}, \|Z\| \lesssim_P T^{1/2}, \|Z\|_{\text{MAX}} \lesssim_P (\log T)^{1/2}, \|\bar{Z} \underline{F}'\| \lesssim_P T^{1/2}, \|\bar{Z} \underline{W}'\| \lesssim_P T^{1/2}, \end{aligned}$$

where $\Sigma_f \in \mathbb{R}^{K \times K}$, $\Sigma_w \in \mathbb{R}^{M \times M}$ are positive-definite matrices with $\lambda_K(\Sigma_f) \gtrsim 1$, $\lambda_M(\Sigma_w) \gtrsim 1$, $\lambda_1(\Sigma_f) \lesssim 1$, and $\lambda_1(\Sigma_w) \lesssim 1$.

Assumption 1 imposes rather weak conditions on the time series behavior of f_t , z_t , and w_t . Since all of them are finite dimensional time series, the imposed inequalities hold if these processes are stationary, strong mixing, and satisfy sufficient moment conditions.

Moreover, Assumption 1 implies that the K left-singular values of F neither vanish nor explode. Therefore, it is the factor loadings that dictate the strength of factors in our setting. This is without loss of generality because F can always be normalized to satisfy this condition.

Next, we assume

Assumption 2. *The $N \times K$ factor loading matrix β satisfies*

$$\|\beta\|_{\text{MAX}} \lesssim 1, \quad \lambda_K(\beta'_{[I_0]} \beta_{[I_0]}) \gtrsim N_0,$$

¹Otherwise, we can define $\tilde{f}_t = f_t - E(f_t w_t') E(w_t w_t')^{-1} w_t$ and $\tilde{\beta}_w = \beta_w + \beta E(f_t w_t') E(w_t w_t')^{-1}$, then $E(\tilde{f}_t w_t') = 0$ and x_t satisfies a similar equation to (1): $x_t = \beta \tilde{f}_t + \tilde{\beta}_w w_t + u_t$.

for some index set $I_0 \subset [N]$, where $N_0 = |I_0| \rightarrow \infty$.

Assumption 2 implies that there exists a subset, I_0 , of predictors within which all latent factors are pervasive. This is a much weaker condition than requiring factors to be pervasive in the set of all predictors, in which case $\lambda_1(\beta'\beta) \asymp \dots \asymp \lambda_K(\beta'\beta) \asymp N$. In contrast, Assumption 2 allows for distinct growth rates for these eigenvalues, in that no requirement is imposed on $\beta_{[I_0^c]}$. Moreover, these eigenvalues can grow at a slower rate than N , since N_0/N is allowed to vanish very rapidly. We will make precise statement about the relative magnitudes of these quantities when it comes to our asymptotic results.

Since the number of factors, K , is assumed finite, even if each factor is pervasive in some separate (and potentially non-overlapping) index set, it is possible to construct a common index set I_0 within which all factors are pervasive.² Assumption 2, nevertheless, rules out a somewhat extreme case where all entires of β are uniformly vanishing, i.e., $\sup_{I, |I| \rightarrow \infty} |I|^{-1} \lambda_K \left(\beta'_{[I]} \beta_{[I]} \right) = o_P(1)$, to the extent that the desired subset I_0 does not exist.

Next, we need the following moment conditions on U .

Assumption 3. *The idiosyncratic component U satisfies:*

$$\|U\|_{\text{MAX}} \lesssim_P (\log T)^{1/2} + (\log N)^{1/2}.$$

In addition, for any given non-random subset $I \subset [N]$,

$$\|U_{[I]}\| \lesssim_P |I|^{1/2} + T^{1/2}.$$

Assumption 3 imposes restrictions on the time-series dependence and heteroskedasticity of u_t . The first inequality is a direct result of a large deviation theorem, see, e.g., Fan et al. (2011). The second inequality can be shown by random matrix theory, see Bai and Silverstein (2009), provided that u_t is i.i.d. both in time and in the cross-section. While it is tempting to impose a stronger inequality that bounds $\sup_{I \subset [N]} \|U_{[I]}\|$ uniformly over all index sets of a given size $|I|$, the rate $|I|^{1/2} + T^{1/2}$ we desire may not hold. In fact, assuming $|I|$ is small, Cai et al. (2021) establish a uniform bound that differs from our non-uniform rate only by a log factor. When $|I|$ is large, the result on uniform bounds no longer exists to the best of our knowledge. We thereby avoid making any assumption on uniform bound over all index sets.

For the same reason, we make the following moment conditions with any given non-random set I . The conditions should hold under weak dependences among U , F , W , and β .

²To see a concrete example, suppose that β has a block diagonal structure, such that its k th column β_k is supported on an index set J_k , and the intersection of all J_k s is empty. Suppose the non-zero entries of β follow standard normal. Then we can find $k^* := \min_k |J_k|$, and build up I_0 from J_{k^*} (so that $|I_0| \geq |J_{k^*}|$) by arbitrarily adding $|J_{k^*}|$ number of predictors from each $J_k, k = 1, 2, \dots, K, k \neq k^*$. We can take a union of all such subsets of J_k . The resulting index set I_0 contains $K \times |J_{k^*}|$ number of predictors, and all factors are pervasive within this common set.

Assumption 4. For any non-random subset $I \subset [N]$, the factor loading $\beta_{[I]}$, and the idiosyncratic error $U_{[I]}$ satisfy the following conditions:

$$\begin{aligned} (i) \quad & \left\| \underline{U}_{[I]} A' \right\| \lesssim_P |I|^{1/2} T^{1/2}, \left\| \underline{U}_{[I]} A' \right\|_{\text{MAX}} \lesssim_P (\log N)^{1/2} T^{1/2}, \\ (ii) \quad & \left\| \beta'_{[I]} U_{[I]} \right\| \lesssim_P |I|^{1/2} T^{1/2}, \left\| \beta'_{[I]} U_{[I]} \right\|_{\text{MAX}} \lesssim_P |I|^{1/2} (\log T)^{1/2}, \left\| \beta'_{[I]} \underline{U}_{[I]} A' \right\| \lesssim_P |I|^{1/2} T^{1/2}, \\ (iii) \quad & \left\| (u_T)'_{[I]} \underline{U}_{[I]} A' \right\| \lesssim_P |I| + |I|^{1/2} T^{1/2}, \left\| \beta'_{[I]} (u_T)_{[I]} \right\| \lesssim_P |I|^{1/2}, \end{aligned}$$

where A is either $\underline{F}$, $\underline{W}$ or $\overline{Z}$.

The ℓ_2 -norm bounds in Assumption 4(i) and (ii) are results of Assumptions D, F2, F3 of Bai (2003) when $I = [N]$, and Assumption 4(iii) is implied by Assumptions A, E1, F1 and C3 in Bai (2003), except that here we impose a stronger version which holds for any non-random subset $I \subset [N]$. The MAX-norm results can be shown by some large deviation theorem as in Fan et al. (2011).

Assumptions 2 and 3 are the key identification conditions of the weak factor model we consider. It is helpful to compare these conditions with those spelled out by Chamberlain and Rothschild (1983). We do not require that u_t is stationary, but for the sake of comparison here, we assume that the covariance matrix of u_t exists, denoted by Σ_u and that $\beta_w = 0$. By model setup (1), we have $\Sigma := \text{Cov}(x_t) = \beta \Sigma_f \beta' + \Sigma_u$. Chamberlain and Rothschild (1983) show that the model is identified if $\|\Sigma_u\| \lesssim 1$ and $\lambda_K(\beta' \beta) \gtrsim N$, which guarantees the separation of the common and idiosyncratic components. Bai (2003) provides an alternative set of conditions (Assumption C therein) on the time-series and cross-sectional dependence of the idiosyncratic components that ensure the consistency of PCA, but also in the case of pervasive factors.

In fact, PCA can separate the factor and idiosyncratic components from the sample covariance matrix under much weaker conditions. To see this, note that from (1) and $\beta_w = 0$, we have $XX' = \beta F F' \beta' + U U' + \beta F U' + U F' \beta'$. Using random matrix theory from Bai and Silverstein (2009), $\lambda_1(U U') \lesssim_P T + N$, if u_t is i.i.d. with $\|\Sigma_u\| \lesssim 1$. Since $T \lambda_K(\beta' \beta) \asymp_P \lambda_K(\beta F F' \beta')$ and because of the weak dependence between U and F as in Assumption 4, the eigenvalues corresponding to the factor component $\beta F F' \beta'$ dominate the three remainder terms that are related to the idiosyncratic component U asymptotically, if $(T + N)/(T \lambda_K(\beta' \beta)) \rightarrow 0$, enabling the factor components to be identified from XX' . Wang and Fan (2017) and Bai and Ng (2021) study the setting $N/(T \lambda_K(\beta' \beta)) \rightarrow 0$, in which case PCA remains consistent despite the fact that factor exposures are not pervasive. Wang and Fan (2017) also study the borderline case $N \asymp T \lambda_K(\beta' \beta)$, and document a bias term in the estimated eigenvalues and eigenvectors associated with factors.

In this paper, we consider an even weaker factor setting in which $N/(T \lambda_K(\beta' \beta))$ may diverge. In this case, PCA generally fails to recover the underlying factors (unless errors are homoscedastic). We will require, instead, the existence of a subset $I_0 \subset [N_0]$, for which $|I_0|/(T \lambda_K(\beta'_{[I_0]} \beta_{[I_0]})) \rightarrow 0$, to ensure the identification of factors on this subset.³ In what follows, we introduce our methodology

³The aforementioned settings all require $\lambda_K(\beta' \beta) \rightarrow \infty$, in contrast with the extremely weak factor model that

to deal with this case.

2.3 Prediction via Supervised Principal Components

One potential solution to the weak factor problem was proposed by [Bair and Tibshirani \(2004\)](#), namely, supervised principal component analysis. Their proposal is to locate a subset, $\widehat{I}$, of predictors via marginal screening, keeping only those that have nontrivial exposure to the prediction target, before applying PCA. Intuitively, this procedure reduces the total number of predictors from N to $|\widehat{I}|$, while under certain assumptions it also guarantees that this subset of predictors has a strong factor structure, i.e., $\lambda_{\min}(\beta'_{[\widehat{I}]} \beta_{[\widehat{I}]}) \asymp |\widehat{I}|$. As a result, applying PCA on this subset leads to consistent recovery of factors.

We use a simple one factor example to illustrate the procedure, before explaining its caveats with the general multi-factor case. To illustrate the idea, we consider the case in which $D = K = 1$, $\alpha_w = 0$, and $\beta_w = 0$. We select a subset $\widehat{I}$ that satisfies:

$$\widehat{I} = \left\{ i \mid |T^{-1} | \underline{X}_{[i]} \overline{Y}' | \geq c \right\}, \quad (3)$$

where c is some threshold. Therefore, we keep predictors that covary sufficiently strongly (positively or negatively) with the target. This step involves a single tuning parameter, c , that effectively determines how many predictors we use to extract the factor. The fact that $\widehat{I}$ incorporates information from the target reflects the distinctive nature of a supervised procedure. Given the existence of I_0 by [Assumption 2](#), there exists a choice of c such that predictors within the set $\widehat{I}$ have a strong factor structure. The rest of the procedure is a straightforward application of the principal component regression for prediction. Specifically, we apply PCA to extract factors $\{\widehat{f}_t\}_{t=1}^{T-h}$ from $\underline{X}_{[\widehat{I}]}$, which can be written as $\widehat{f}_t = \widehat{\zeta}' x_t$ for some loading matrix $\widehat{\zeta}$, then obtain $\widehat{\alpha}$ by regressing $\{y_t\}_{t=1+h}^T$ onto $\{\widehat{f}_t\}_{t=1}^{T-h}$ based on the predictive model [\(2\)](#). The resulting predictor for y_{T+h} is therefore given by: $\widehat{y}_{T+h} = \widehat{\alpha} \widehat{f}_T = \widehat{\alpha} \widehat{\zeta}' x_T$.

[Bair et al. \(2006\)](#)'s proposal proceeds in the same way when it comes to multiple factors, with the only exception that multiple factors are extracted in the PCA step. Yet, to ensure that marginal screening remains valid in the multi-factor setting, they assume that predictors are marginally correlated with the target *if and only if* they belong to a *uniquely* determined subset I_0 , outside which predictors are assumed to have zero correlations with the prediction target, i.e., they are pure noise for prediction purpose. Given this condition, they show marginal screening can consistently recover I_0 , and all factors can thereby be extracted altogether with a single pass of PCA to this subset of predictors.

In contrast, we assume the existence of a set I_0 within which predictors have a strong factor

imposes $\lambda_K(\beta' \beta) \lesssim 1$. As such, eigenvalues of factors and idiosyncratic components do not diverge as dimension increases. While [Onatski \(2009\)](#) and [Onatski \(2010\)](#) develop tests for the number of factors, [Onatski \(2012\)](#) shows that factors cannot be consistently recovered in this regime.

structure, yet we do not make any assumptions on the correlation between the target and predictors outside this set I_0 , nor on the strength of their factor structure. As a result, I_0 under our Assumption 2 needs not be unique, and we will show that the validity of the prediction procedure does not rely on consistent recovery of any pre-determined set I_0 . More importantly, since marginal screening is based on marginal covariances between $\bar{Y}$ and $\underline{X}$, in a multi-factor model the condition that marginal screening can recover a subset within which all factors are pervasive (even if such a subset is uniquely defined as in Bair et al. (2006)) is rather strong. On the one hand, marginal screening can be misguided by the correlation induced by a strong factor to the extent that weak factors after screening remain unidentifiable. On the other hand, predictors eliminated by marginal screening can be instrumental or even essential for prediction. We illustrate these points using examples of two-factor models below.

EXAMPLE 1: Suppose x_t and y_t satisfy the following dynamics:

$$x_t = \left[\begin{array}{c|c} \beta_{11} & \beta_{12} \\ \hline \beta_{21} & 0 \end{array} \right] f_t + u_t, \quad y_{t+h} = \begin{bmatrix} 1 & 1 \end{bmatrix} f_t, \quad (4)$$

where β_{11} and β_{12} are $N_0 \times 1$ vectors, β_{21} is an $(N - N_0) \times 1$ vector, satisfying $\|\beta_{12}\| \asymp N_0^{1/2}$ and $\|\beta_{21}\| \asymp (N - N_0)^{1/2}$, and N_0 is small relative to N .

In this example, the first factor is strong (all predictors are exposed to it) while the second factor is weak, since most exposures to it are zero. In addition, the target variable y is correlated with both factors and hence potentially with all predictors. As a result, the screening step described above may not eliminate any predictors: all predictors may correlate with the target (through the first factor). But because the second factor is weak, a single pass of PCA, extracting two factors from the entire universe of predictors, would fail to recover it: we can show that $\lambda_{\min}(\beta' \beta) \leq \|\beta_{12}\|^2 \lesssim N_0$, so that PCA would not recover the second factor consistently if $N/(N_0 T)$ does not vanish.

The issue highlighted with this example is that the (single) screening step does not eliminate any predictors, because their correlations with the target are (at least partially) induced by their exposure to the strong factor, and therefore PCA after screening cannot recover the weak factor. The assumptions proposed by Bair et al. (2006) rule this case out, but we can clearly locate an index set I_0 (say, top N_0 predictors), within which both factors are strong. In other words, our assumptions can accommodate this case.

We provide next another example, that shows that in some situations screening can eliminate *too many* predictors, making a strong factor model become weak or even rank-deficient.

EXAMPLE 2: Suppose x_t and y_t satisfy the following dynamics:

$$x_t = \left[\begin{array}{c|c} \beta_{11} & \beta_{11} \\ \hline 0 & \beta_{22} \end{array} \right] f_t + u_t, \quad y_{t+h} = \begin{bmatrix} 1 & 0 \end{bmatrix} f_t, \quad (5)$$

where β_{11} and β_{22} are $N/2 \times 1$ non-zero vectors satisfying $\|\beta_{11}\| \asymp \|\beta_{22}\| \asymp \sqrt{N}$ and f_{1t} and f_{2t} are uncorrelated.

In this example, there are two equal-sized groups of predictors, so that β is full-rank and both factors are strong and that I_0 can be the entire set $[N]$ (therefore, a standard PCA procedure applied to all predictors will consistently recover both factors). But two features of this model will make supervised PCA fail, if the selection step based on marginal correlations is applied only once (as in the original procedure by [Bair et al. \(2006\)](#)). First, y_{t+h} is uncorrelated with the second half of predictors (since only the first group is useful for prediction). Second, the exposure of the first half of predictors to the first and second factors are the same (both equal to β_{11}).

After the screening step the second group of predictors would be eliminated, because they do not marginally correlate with y_{t+h} . But the remaining predictors (the first half) have perfectly correlated exposures to both factors, so that only one factor, $f_{1t} + f_{2t}$, can be recovered by PCA. Therefore, the one-step supervised PCA of [Bair et al. \(2006\)](#) would fail to recover the factor space consistently, resulting in inconsistent prediction. This example highlights an important point that marginally uncorrelated predictors (the second half) could be essential in recovering the factor space. Eliminating such predictors may lead to inconsistency in prediction.

Both examples demonstrate the failure of a one-step supervised PCA procedure in a general multi-factor setting. Such data generating processes are excluded by the model assumptions in [Bair et al. \(2006\)](#), whereas we do not rule them out. We thus propose below a new and more complete version of the supervised PCA (SPCA) procedure that can accommodate such cases.

2.4 Iterative Screening and Projection

To resolve the issue of weak factors in a general multi-factor setting, we propose a multi-step procedure that iteratively conducts selection and projection. The projection step eliminates the influence of the estimated factor, which ensures the success of the screening steps that occur over the following iterations. More specifically, a screening step can help identify one strong factor from a selected subset of predictors. Once we have recovered this factor, we project *all* predictors x_t (not just those selected at the first step) and y_{t+h} onto this factor, so that their residuals will not be correlated with this factor. Then we can repeat the same selection procedure with these residuals. This approach enables a continued discovery of factors, and guarantees that each new factor is orthogonal to the

estimated factors in the previous steps, similar to the standard PCA.

It is straightforward to verify that this iterative screening and projection approach successfully addresses the issues with the aforementioned examples. Consider first Example 1. In this case, the first screening does not rule out any predictor, and the first PC will recover the strong factor f_1 ; after projecting both X and y onto f_1 , the residuals for the first N_0 predictors still load on f_2 , whereas the remaining $N - N_0$ predictors should have zero correlation with the residuals of y . Therefore, a second screening will eliminate these predictors, paving the way for PCA to recover the second factor f_2 based on the residuals of the first N_0 predictors. Similarly, for Example 2, the first screening step eliminates the second half of the predictors, so that the first pass of PCA will recover the only factor left over in the remaining predictors, namely, $f_1 + f_2$. The residuals of the first half of predictors consist of pure noise after the projection step, whereas the residuals of the second half of predictors are spanned by $f_1 - f_2$, which a second PCA step will recover. Therefore, the iterated supervised PCA will recover the entire factor space. This example illustrates that marginal screening can succeed as long as iteration and projection are also employed.

Formally, we present our algorithm for the general model given by (1) and (2):

Algorithm 1 (Prediction via SPCA).

Inputs: $\bar{Y}$, $\underline{X}$, $\underline{W}$, x_T , and w_T . *Initialization:* $Y_{(1)} := \bar{Y}\mathbb{M}_{\underline{W}'}$, $X_{(1)} := \underline{X}\mathbb{M}_{\underline{W}'}$.

S1. For $k = 1, 2, \dots$ iterate the following steps using $X_{(k)}$ and $Y_{(k)}$:

- a. Select an appropriate subset $\hat{I}_k \subset [N]$ via marginal screening.*
- b. Estimate the k th factor $\hat{\underline{F}}_{(k)} = \hat{\zeta}_{(k)}' (X_{(k)})_{[\hat{I}_k]}$ via SVD, where $\hat{\zeta}_{(k)}$ is the first left singular vector of $(X_{(k)})_{[\hat{I}_k]}$. $\hat{\underline{F}}_{(k)}$ can also be rewritten as $\hat{\underline{F}}_{(k)} = \hat{\zeta}_{(k)}' \underline{X}\mathbb{M}_{\underline{W}'}$, where $\hat{\zeta}_{(k)} = \left(\mathbb{I}_N - \sum_{i=1}^{k-1} \hat{\beta}_{(i)} \hat{\zeta}_{(i)}' \right)'_{[\hat{I}_k]}$ $\hat{\zeta}_{(k)}$ is constructed recursively using $\hat{\beta}_{(k-1)}$ (defined in c.).*
- c. Estimate the coefficients $\hat{\alpha}_{(k)} = Y_{(k)} \hat{\underline{F}}_{(k)}' (\hat{\underline{F}}_{(k)} \hat{\underline{F}}_{(k)}')^{-1}$ and $\hat{\beta}_{(k)} = X_{(k)} \hat{\underline{F}}_{(k)}' (\hat{\underline{F}}_{(k)} \hat{\underline{F}}_{(k)}')^{-1}$.*
- d. Obtain residuals $Y_{(k+1)} = Y_{(k)} - \hat{\alpha}_{(k)} \hat{\underline{F}}_{(k)}$ and $X_{(k+1)} = X_{(k)} - \hat{\beta}_{(k)} \hat{\underline{F}}_{(k)}$.*

Stop at $k = \hat{K}$, where $\hat{K}$ is chosen based on some proper stopping rule.

S2. Obtain $\hat{f}_T = \hat{\zeta}' (x_T - \hat{\beta}_w w_T)$, where $\hat{\zeta} := (\hat{\zeta}_{(1)}, \dots, \hat{\zeta}_{(\hat{K})})$ and $\hat{\beta}_w = \underline{X}\underline{W}'(\underline{W}\underline{W}')^{-1}$, and the prediction $\hat{y}_{T+h} = \hat{\alpha} \hat{f}_T + \hat{\alpha}_w w_T = \hat{\gamma} x_T + (\hat{\alpha}_w - \hat{\gamma} \hat{\beta}_w) w_T$, where $\hat{\alpha} := (\hat{\alpha}_{(1)}, \hat{\alpha}_{(2)}, \dots, \hat{\alpha}_{(\hat{K})})$, $\hat{\gamma} = \hat{\alpha} \hat{\zeta}'$, and $\hat{\alpha}_w = \bar{Y}\underline{W}'(\underline{W}\underline{W}')^{-1}$.

Outputs: the prediction $\hat{y}_{T+h}$, the factors $\hat{\underline{F}} := (\hat{\underline{F}}_{(1)}', \dots, \hat{\underline{F}}_{(\hat{K})}')'$, their loadings, $\hat{\beta} := (\hat{\beta}_{(1)}, \dots, \hat{\beta}_{(\hat{K})})$, and the coefficient estimates $\hat{\alpha}$, $\hat{\zeta}$, $\hat{\alpha}_w$, $\hat{\beta}_w$, and $\hat{\gamma}$.

We discuss the details of the algorithm below.

Step S1. of Algorithm 1 requires an appropriate choice of $\hat{I}_k$ and a stopping rule. One possible choice for $\hat{I}_k$ is:⁴

$$\hat{I}_k = \left\{ i \mid T^{-1} \left\| (X_{(k)})_{[i]} Y'_{(k)} \right\|_{\text{MAX}} \geq \hat{c}_{qN}^{(k)} \right\},$$

where $\hat{c}_{qN}^{(k)}$ is the $(1-q)th$ -quantile of $\left\{ T^{-1} \left\| (X_{(k)})_{[i]} Y'_{(k)} \right\|_{\text{MAX}} \right\}_{i=1, \dots, N}$. (6)

The reason we suggest using the top qN predictors based on the magnitude of the covariances between $X_{(k)}$ and $Y_{(k)}$ is that the factor estimates tend to be more stable and less sensitive to this tuning parameter q , compared to a conventional hard threshold parameter adopted in a marginal screening procedure. Moreover, at each step, a subset of a *fixed* number of predictors are selected, which substantially simplifies the notation and the proof.

Correspondingly, the algorithm terminates as soon as

$$\hat{c}_{qN}^{(k+1)} < c, \quad \text{for some threshold } c. \quad (7)$$

Thus, the resulting number of factors is set as $\hat{K} = k$. As a result, the tuning parameter, c , effectively determines the number of factors extracted out of our procedure.

For any given tuning parameters, q and c , we select predictors that have predictive power for (at least one variable in) y_{t+h} at each stage of the iteration. With a good choice of tuning parameters, q and c , the iteration stops as soon as most of the rows of the projected residuals of predictors appear uncorrelated with the projected residuals of y_{t+h} , which implies that the factors left over, if any, are uncorrelated with y_{t+h} .

The last step of the algorithm needs more explanations. Step S1. provides a set of factor estimates, $\hat{F}$, on the basis of $\bar{Y}$ and $\underline{X}$. Moreover, a time series regression of $\bar{Y}$ on $\hat{F}$ and $\underline{W}$ yields an estimator of α_w (coefficient defined in (2)). That is, $\hat{\alpha}_w = \bar{Y} \mathbb{M}_{\hat{F}}' \underline{W}' \left(\underline{W} \mathbb{M}_{\hat{F}}' \underline{W}' \right)^{-1} = \bar{Y} \underline{W}' (\underline{W} \underline{W}')^{-1}$, since $\mathbb{M}_{\hat{F}}' \underline{W}' = \underline{W}'$ by construction, which explains the formula for $\hat{\alpha}_w$ in Step S2.. Finally, with $\hat{\alpha}$, $\hat{\alpha}_w$, and $\hat{f}_T$, it is sufficient to construct the predicted value of y_{T+h} by combining $\hat{\alpha} \hat{f}_T$ with $\hat{\alpha}_w w_T$, which yields the final prediction formula for $\hat{y}_{T+h}$, a projection on observables, x_T and w_T .

3 Asymptotic Theory

We now examine the asymptotic properties of SPCA. The analysis is more involved than those of [Bair et al. \(2006\)](#) because of the iterative nature of our new SPCA procedure and the general weak factor setting we consider.

⁴Using covariance for screening allows us to replace all $Y_{(k)}$ in the definition of $\hat{I}_k$ and Algorithm 1 by $Y_{(1)}$, that is, only the projection of $X_{(k)}$ is needed, because this replacement would not affect the covariance between $Y_{(k)}$ and $X_{(k)}$. We use this fact in the proofs, which simplifies the notation. We can also use correlation instead of covariance in constructing $\hat{I}_k$, which does not affect the asymptotic analysis. That said, we find correlation screening performs better in finite samples when the scale of the predictors differs.

3.1 Consistency in Prediction

To establish the consistency of SPCA for prediction, we first investigate the consistency of factor estimation. In the strong factor case, e.g., [Stock and Watson \(2002\)](#), all factors are recovered consistently via PCA, which is a prerequisite for the consistency of prediction. In our setup of weak factors, we show that the consistency of prediction only relies on consistent recovery of factors that are relevant for the prediction target.

Recall that in [Algorithm 1](#), we denote the selected subsets in the SPCA procedure as $\hat{I}_k$, $k = 1, 2, \dots$. We now construct their population counterparts iteratively, for any given choice of c and q . This step is critical to characterize the exact factor space recovered by SPCA. For simplicity in notation and without loss of generality, we consider the case $\Sigma_f = \mathbb{I}_K$ here, because in the general case, we can simply replace β and α by $\beta^* = \beta \Sigma_f^{1/2}$ and $\alpha^* = \alpha \Sigma_f^{1/2}$ in the following construction.

In detail, we start with $a_i^{(1)} := \|\beta_{[i]} \alpha'\|_{\text{MAX}}$ and define $I_1 := \{i | a_i^{(1)} \geq c_{qN}^{(1)}\}$, where $c_{qN}^{(1)}$ is the $[qN]$ th largest value in $\{a_i^{(1)}\}_{i=1, \dots, N}$. Then, we denote the largest singular value of $\beta_{(1)} := \beta_{[I_1]}$ by $\lambda_{(1)}^{1/2}$ and the corresponding left and right singular vectors by $\varsigma_{(1)}$ and $b_{(1)}$. For $k > 1$, we obtain $a_i^{(k)} := \left\| \beta_{[i]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}}$, $I_k := \{i | a_i^{(k)} \geq c_{qN}^{(k)}\}$, and $\lambda_{(k)}^{1/2}$, $\varsigma_{(k)}$, $b_{(k)}$ are the leading singular value, left and right singular vectors of $\beta_{(k)} := \beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}}$. This procedure is stopped at step $\tilde{K}$ (for some $\tilde{K}$ that is not necessarily equal to K or $\hat{K}$) if $c_{qN}^{(\tilde{K}+1)} < c$. In a nutshell, I_k 's are what we will select if we do SPCA directly on $\beta \in \mathbb{R}^{N \times K}$ and $\alpha \in \mathbb{R}^{D \times K}$ and they are deterministically defined by $\alpha, \beta, \Sigma_f, c, q$, and N , whereas $\hat{I}_k$'s are random, obtained by SPCA on $\underline{X} \in \mathbb{R}^{N \times T}$ and $\bar{Y} \in \mathbb{R}^{D \times T}$.

To ensure that the singular vectors $b_{(j)}$'s are well defined and identifiable, we need that the top two singular values of $\beta_{(k)}$ are distinct at each stage k . We also need distinct values of $c_{qN}^{(k)}$ to ensure that I_k 's are identifiable. More precisely, we say that two sequences of variables a_N and b_N are asymptotically distinct if there exists a constant $\delta > 0$ such that $a_N \leq (1 + \delta)^{-1} b_N$. In light of the above discussion, we make the following assumption:

Assumption 5. *For any given k , the following three pairs of sequences of variables, $\sigma_1(\beta_{(k)})$ and $\sigma_2(\beta_{(k)})$, $c_{qN}^{(k)}$ and $c_{qN+1}^{(k)}$, and $c_{qN}^{(\tilde{K}+1)}$ and c are asymptotically distinct, as $N \rightarrow \infty$.*

This assumption is rather mild as it only rules out corner cases, despite the fact that this is not very explicit. Excluding such corner cases is common in the literature on high dimensional PCA, see, e.g., [Assumption 2.1 of Wang and Fan \(2017\)](#). We now are ready to present the consistency of the estimated factors by SPCA:

Theorem 1. *Suppose that x_t follows [\(1\)](#) and y_t satisfies [\(2\)](#), and that [Assumptions 1-5](#) hold. If $\log(NT)(N_0^{-1} + T^{-1}) \rightarrow 0$, then for any tuning parameters c and q that satisfy*

$$c \rightarrow 0, \quad c^{-1}(\log NT)^{1/2}(q^{-1/2}N^{-1/2} + T^{-1/2}) \rightarrow 0, \quad qN/N_0 \rightarrow 0, \quad (8)$$

we have $\tilde{K} \leq K$, $P(\hat{I}_k = I_k) \rightarrow 1$, for any $1 \leq k \leq \tilde{K}$, and $P(\hat{K} = \tilde{K}) \rightarrow 1$. Moreover, the factors recovered by SPCA are consistent. That is, for any $1 \leq k \leq \tilde{K}$,

$$\left\| \hat{\underline{F}}_{(k)} \right\|^{-1} \left\| \hat{\underline{F}}_{(k)} - \underline{F}_{(k)} \mathbb{P}_{\underline{F}'} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}. \quad (9)$$

We make a few observations regarding this result. First, the assumptions in Theorem 1 do not guarantee a consistent estimate of the number of factors, K , because the SPCA procedure cannot guarantee to recover factors that are uninformative about y . At the same time, the factors recovered by SPCA are not necessarily useful for prediction, because it is possible that some strong factors with no predictive power are also recovered by SPCA. Ultimately, the factor space recoverable is determined by β , α , Σ_f , c , q , and N . For this reason, we have consistency of factor estimates up to the first $\tilde{K}$ factors. Moreover, $\hat{K}$ is a consistent estimator of $\tilde{K}$, which we prove satisfies $\tilde{K} \leq K$. That is, SPCA omits $K - \tilde{K}$ factors. Also, the inequality (9) has a clear geometric interpretation. The left-hand-side is exactly equal to $\sin(\hat{\Theta}_{(k)})$, where $\hat{\Theta}_{(k)}$ is the angle between the estimated factor at each stage k and the factor space spanned by the true factors, $\mathbb{P}_{\underline{F}'}$. (9) shows that this angle vanishes asymptotically.

Second, with respect to the tuning parameters, the condition (8) implies that $c \rightarrow 0$, $c\sqrt{T} \rightarrow \infty$, and $c\sqrt{qN} \rightarrow \infty$. On the one hand, the threshold c needs be sufficiently small so that the iteration procedure continues until selected predictors have asymptotically vanishing predictive power; on the other hand, c needs be large enough that dominates error in the covariance estimates from the screening step. The estimation error consists of the usual error in the construction of the sample covariances between $X_{(1)}$ and $Y_{(1)}$, which introduces an error of order $T^{-1/2}$, as well as the construction of residuals in the projection step, $X_{(k)}$ and $Y_{(k)}$, for $k > 1$, as soon as multiple factors are involved (i.e., $\tilde{K} > 1$). As we show next, the factor estimation error is of order $(qN)^{-1/2} + T^{-1}$, which pollutes the residuals and hence affects screening. Taking these two points into consideration, the choice of c needs dominate $T^{-1/2} + (qN)^{-1/2}$. In terms of q , it appears that the maximal number of selected predictors, $\lfloor qN \rfloor$, allowed for should be of the same order as N_0 . Nevertheless, since N_0 given by Assumption 2 is not precisely defined, in the sense that the assumption holds if N_0 is scaled by any non-zero constant, we require $qN/N_0 \rightarrow 0$ to ensure that the scaling constant of N_0 does not matter for the choice of q and that the selected $\lfloor qN \rfloor$ predictors are within the subset of N_0 predictors that guarantee a strong factor structure.

Third, the estimation error of factors are bounded from the above by $q^{-1/2} N^{-1/2} + T^{-1}$. Recall that in the strong factor case, the factor space can be recovered at the rate of $N^{-1/2} + T^{-1}$, see, e.g., Bai (2003). In our result, qN plays the same role as N in the strong factor case. Nevertheless, our Assumption 2 does not require all factors to have the same strength. It is possible that some factors could be recovered with a higher convergence rate, should we select a different number of predictors for each factor based on its strength. In fact, an alternative choice of $\hat{I}_k$ based on (3) allows different numbers of predictors to be selected at each stage, since the threshold itself is a fixed

level. While this approach may achieve a faster rate for relatively stronger factors, the prediction error rate is ultimately determined by the estimation error of the weakest factor. Yet, we find that the approach based on (6) offers more stable prediction out of sample, whereas prediction based on (3) can be sensitive to the tuning parameters. Given that our ultimate goal is about prediction rather than factor recovery, we prefer a more stable procedure and thereby focus our analysis on the former approach.

With no relevant factors omitted, our prediction $\hat{y}_{T+h}$ is consistent, as we show next.

Theorem 2. *Under the same assumptions as in Theorem 1, we have $\hat{\alpha}_w - \alpha_w \xrightarrow{P} 0$, $\|\hat{\gamma}\beta - \alpha\| \xrightarrow{P} 0$, and consequently, $\hat{y}_{T+h} \xrightarrow{P} E_T(y_{T+h}) = \alpha f_T + \alpha_w w_T$.*

Theorem 2 first analyzes the parameter estimation “error” measured as $\hat{\alpha}_w - \alpha_w$ and $\hat{\gamma}\beta - \alpha$. The reason the latter quantity matters is that there exists a matrix H such that $\hat{\gamma}\beta = \hat{\alpha}H$. In other words, the first statement of the theorem implies that we can consistently estimate α , up to a matrix H . This extra adjustment matrix H exists due to the fundamental indeterminacy of latent factor models. In fact, we can define $H \in \mathbb{R}^{\hat{K} \times K}$ as $\hat{\zeta}'\beta$, where $\hat{\zeta}$ is given by Algorithm 1. Then, it is straightforward to see from the definition of $\hat{\gamma}$ that

$$\hat{\gamma}\beta = \hat{\alpha}H, \quad \text{so that by Theorem 2} \quad \|\hat{\alpha}H - \alpha\| = o_P(1). \quad (10)$$

On the other hand, the proof of Theorem 1 also establishes that for $k \leq \tilde{K}$:

$$\left\| \hat{\underline{F}}_{(k)} \right\|^{-1} \left\| \hat{\underline{F}}_{(k)} - h_k \underline{F} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}, \quad (11)$$

where h_k is the k th row of H . Therefore, $\hat{\alpha}\hat{\underline{F}} \stackrel{\text{by(11)}}{\approx} \hat{\alpha}H\underline{F} \stackrel{\text{by(10)}}{\approx} \alpha\underline{F}$, which, together with $\hat{\alpha}_w - \alpha_w = o_P(1)$, leads to the consistency of prediction.

The consistency result in Theorem 2 does not require a full recovery of all factors. In other words, $\hat{K}$ is not necessarily equal to K . On the one hand, factors omitted by SPCA are guaranteed to be uncorrelated with y_{t+h} ; on the other hand, some factors not useful for prediction may be recovered by SPCA. Obviously, missing any uncorrelated factors or having extra useless factors (for prediction purposes) do not affect the consistency of $\hat{y}_{T+h}$.

Moreover, this result does not rely on normally distributed error nor on the assumption that all factors share the same strength with respect to all predictors. The assumption on the relative size of N and T is also quite flexible, in contrast with existing results in the literature in which N cannot grow faster than a certain polynomial rate of T , e.g., Bai and Ng (2021), Huang et al. (2021).

3.2 Recovery of All Factors

In this section we develop the asymptotic distribution of $\hat{y}_{T+h}$ from Algorithm 1. Not surprisingly, the conditions in Theorem 2 are inadequate to guarantee that $\hat{y}_{T+h}$ converges to $E_T(y_{T+h})$ at the

desirable rate $T^{-1/2}$. The major obstacle lies in the recovery of all factors, which we will illustrate with a one-factor example.

EXAMPLE 3: Suppose that x_t follows a single-factor model with sparse β :

$$x_t = \begin{bmatrix} \beta_1 \\ 0 \end{bmatrix} f_t + u_t, \quad y_{t+h} = \alpha f_t + z_{t+h},$$

where β_1 is the first N_0 entries of β with $\|\beta_1\| \asymp N_0^{1/2}$ and $\alpha \asymp T^{-1/2}$.

Recall that we use the sample covariance between x_t and y_{t+h} to screen predictors. Even if y_{t+h} is independent of x_t , their sample covariance can be as large as $T^{-1/2}(\log N)^{1/2}$. Therefore, the threshold c needs be strictly greater than $T^{-1/2}(\log N)^{1/2}$ to control Type I error in screening. However, the signal-to-noise ratio in this example is rather low, i.e., $\alpha \asymp T^{-1/2}$, that is, y_{t+h} is not too different from random noise. Consequently, screening will terminate right away because the covariances between y_{t+h} and x_t are at best of order $T^{-1/2}(\log N)^{1/2} < c$, which in turn leads to no discovery of factors. Our procedure thereby gives $\hat{y}_{T+h} = 0$, which is certainly consistent as the bias $|\mathbb{E}_T(y_{T+h}) - 0| \asymp T^{-1/2}$, but the usual central limit theorem (CLT) fails.

Generally speaking, this issue arises because of the potential failure to recover all factors in the DGP. As long as all factors are found, the bias is negligible and the central limit theorem holds regardless of the magnitude of α . So to go beyond consistency and make valid inference we need a stronger assumption that rules out cases like this, in order to insure against a higher order omitted factor bias that impedes the CLT even if it does not affect consistency. It turns out that as long as $\alpha \in \mathbb{R}^{D \times K}$ satisfies $\lambda_{\min}(\alpha' \alpha) \gtrsim 1$, we can rule out the possibility of missing factors asymptotically. On the one hand, in this case the dimension of target variables, D , must be no smaller than the dimension of the factors, K ; and for each factor there exist at least one target variable in y that is correlated with the factor; together they guarantee that no factors would be omitted. On the other hand, our algorithm will not select more factors than needed asymptotically, because the iteration is terminated as soon as all covariances vanish. With a consistent estimator of the number of factors, we can recover the factor space as well as conduct inference on the prediction targets.

The inference theory on strong factor models also relies on a consistent estimator of the count of (strong) factors, e.g., [Bai and Ng \(2006\)](#). Our assumptions here are substantially weaker than the pervasive factor assumption adopted in the literature. That said, in a finite sample, a perfect recovery of the number of factors may be a stretch. In [Section 3.4](#), we show that our version of the PCA regression is more robust than the procedure of [Stock and Watson \(2002\)](#) with respect to the error due to overestimating the number of factors. We also provide simulation evidence on the finite sample performance of our estimator of the number of factors.

The next theorem summarizes a set of stronger asymptotic results under conditions that guarantee perfect recovery of all factors:

Theorem 3. *Under the same assumptions as Theorem 2, if we further have $\lambda_{\min}(\alpha'\alpha) \gtrsim 1$, then for any tuning parameters c and q in (6) and (7) satisfying*

$$c \rightarrow 0, \quad c^{-1}(\log NT)^{1/2}(q^{-1/2}N^{-1/2} + T^{-1/2}) \rightarrow 0, \quad qN/N_0 \rightarrow 0,$$

we have

(i) $\hat{K}$ defined in Algorithm 1 satisfies: $P(\hat{K} = K) \rightarrow 1$.

(ii) The factor space is consistently recovered in the sense that

$$\left\| \mathbb{P}_{\hat{\underline{F}}} - \mathbb{P}_{\underline{F}'} \right\| = O_P \left(q^{-1/2}N^{-1/2} + T^{-1} \right).$$

(iii) The estimator $\hat{\gamma}$ constructed via Algorithm 1 satisfies

$$\left\| \hat{\gamma}\beta - \alpha - T^{-1}\bar{Z}\underline{F}'\Sigma_f^{-1} \right\| = O_P(q^{-1}N^{-1} + T^{-1}).$$

Theorem 3 extends the strong factor case of Bai and Ng (2002) and Bai (2003). In particular, (i) shows that our procedure can recover the true number of factors asymptotically, which extends Bai and Ng (2002) to the case of weak factors. Combining this result with Theorem 1(i) suggests that $\tilde{K} = K$ under the strengthened set of assumptions. We thereby do not need distinguish $\tilde{K}$ with K below. Our setting is distinct from that of Onatski (2010), and as a result we can also recover the space spanned by weak factors, as shown by (ii). This result also suggests that the convergence rate for factor estimation is of order $(qN)^{1/2} \wedge T$, as opposed to $N^{1/2} \wedge T$ given by Theorem 1 of Bai (2003). (iii) extends the result of Theorem 2, replacing the target α by $\alpha + T^{-1}\bar{Z}\underline{F}'\Sigma_f^{-1}$. Note that the latter is precisely a regression estimator of α if F were observable. (iii) thereby points out that the error due to latent factor estimation is no larger than $O_P(q^{-1}N^{-1} + T^{-1})$.

3.3 Inference on the Prediction Target

In the case without observable regressors w , the prediction error can be written as $\hat{y}_{T+h} - E_T(y_{T+h}) = (\hat{\gamma}\beta - \alpha)f_T + \hat{\gamma}u_T$, where the second term $\hat{\gamma}u_T$ is of order $(qN)^{-1/2}$. In light of Theorem 3(iii), if $q^{-1}N^{-1}T \rightarrow 0$, then the second term is asymptotically negligible (i.e., $o_P(T^{-1/2})$) compared to the first term, $(\hat{\gamma}\beta - \alpha)f_T = T^{-1}\bar{Z}\underline{F}'\Sigma_f^{-1}f_T + O_P(T^{-1})$, in which case we can achieve root- T inference on $E_T(y_{T+h})$. Nevertheless, we strive to achieve a better approximation to the finite sample performance by taking into account both terms of the prediction error altogether, without imposing additional restriction on the relative magnitude of qN and T .

To do so, we impose the following assumption:

Assumption 6. As $N, T \rightarrow \infty$, $T^{-1/2}\bar{Z}\underline{F}'$, $T^{-1/2}\bar{Z}\underline{W}'$, and $(qN)^{-1/2}\Psi u_T$ are jointly asymptotically normally distributed, satisfying:

$$\begin{pmatrix} \text{vec}(T^{-1/2}\bar{Z}\underline{F}') \\ \text{vec}(T^{-1/2}\bar{Z}\underline{W}') \\ (qN)^{-1/2}\Psi u_T \end{pmatrix} \xrightarrow{d} \mathcal{N} \left(\begin{pmatrix} 0 \\ 0 \\ 0 \end{pmatrix}, \Pi = \begin{pmatrix} \Pi_{11} & \Pi_{12} & 0 \\ \Pi'_{12} & \Pi_{22} & 0 \\ 0 & 0 & \Pi_{33} \end{pmatrix} \right),$$

where Ψ is a $K \times N$ matrix whose k th row is equal to $b'_{(k)}\beta'_{[I_k]}(\mathbb{I}_N)_{[I_k]}$ and $b_{(k)}$ is the first right singular vector of $\beta_{(k)} = \beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}}$ as defined in Section 3.1.

Assumption 6 characterizes the joint asymptotic distribution of $\bar{Z}\underline{F}'$, $\bar{Z}\underline{W}'$ and Ψu_T . For the first two components, as the dimensions of these random processes are finite, this CLT is a direct result of a large- T central limit theory for mixing processes. With respect to Ψu_T , its large- N asymptotic distribution is assumed normal, asymptotically independent of the distribution of the other two components. This holds trivially if u_{iT} 's are cross-sectionally i.i.d., independent of z_t , w_t , and f_t for $t < T$, so that the k th row of Ψu_T , $b'_{(k)}\beta'_{[I_k]}(u_T)_{[I_k]}$, is a weighted average of u_{iT} for $i \in I_k$. The convergence rate $(qN)^{1/2}$ for Ψu_T arises naturally because $|I_k| = qN$.

Before we present the CLT next, we need define a $K \times K$ matrix $\Omega = (\omega_1, \dots, \omega_K)$ with $\omega_1 = e_1$ and $\omega_k = e_k - \sum_{i=1}^{k-1} \lambda_{(i)}^{-1} b'_{(k)}\beta'_{[I_k]}\beta_{[I_k]}b_{(i)}\omega_i$, where e_k is a K -dimensional unit vector with 1 on the k th entry and 0 elsewhere.

Theorem 4. Suppose the same assumptions as in Theorem 3 hold. If in addition, Assumption 6 holds, we have

$$\Phi^{-1/2}(\hat{y}_{T+h} - E_T(y_{T+h})) \xrightarrow{d} \mathcal{N}(0, \mathbb{I}_D),$$

where $\Phi = T^{-1}\Phi_1 + q^{-1}N^{-1}\Phi_2$ and Φ_1 and Φ_2 are given by

$$\begin{aligned} \Phi_1 &= \left((f'_T, w'_T) \Sigma_{f,w}^{-1} \otimes \mathbb{I}_D \right) \begin{pmatrix} \Pi_{11} & \Pi_{12} \\ \Pi'_{12} & \Pi_{22} \end{pmatrix} \left(\Sigma_{f,w}^{-1} (f'_T, w'_T)' \otimes \mathbb{I}_D \right), \\ \Phi_2 &= \alpha B(\Lambda/qN)^{-1} \Omega' \Pi_{33} \Omega (\Lambda/qN)^{-1} B' \alpha', \end{aligned}$$

Π_{ij} is specified by Assumption 6, $\Sigma_{f,w} = \text{diag}(\Sigma_f, \Sigma_w)$, $\Lambda = \text{diag}(\lambda_{(1)}, \dots, \lambda_{(K)})$, and B is a $K \times K$ matrix whose k th column is given by $b_{(k)}$, where $\lambda_{(k)}^{1/2}$ is the largest singular value of $\beta_{(k)}$ and $b_{(k)}$ is the corresponding right singular vector as defined in Section 3.1.

The convergence rate of $\hat{y}_{T+h}$ depends on the relative magnitudes of T and qN . For inference, we need construct estimators for each component of Φ_1 and Φ_2 . Estimating Φ_1 is straightforward based on its sample analog, constructed from the outputs of Algorithm 1. Estimating Φ_2 is more involved, in that Π_{33} depends on the large covariance matrix of u_T . We leave the details to the appendix.

Algorithm 1 (Step S2.) makes predictions by exploiting the projection of y_{T+h} onto x_T and w_T , with loadings given by γ and $\alpha_w - \gamma\beta_w$. This is convenient and easily extendable out of sample, as both x_T and w_T are directly observable, unlike latent factors. The next section investigates potential issues with plain PCA and PLS, as well as an alternative algorithm based on [Stock and Watson \(2002\)](#), which does not involve the projection parameter γ .

3.4 Alternative Procedures

In this section, we at first discuss the failure of PCA and PLS in the presence of weak factors. To illustrate the issue, it is sufficient to consider a one-factor model example:

EXAMPLE 4: Suppose that x_t follows a single-factor model with sparse β :

$$x_t = \begin{bmatrix} \beta_1 \\ 0 \end{bmatrix} f_t + u_t, \quad y_{t+h} = \alpha f_t,$$

where β_1 is the first N_0 entries of β with $\|\beta_1\| \asymp N_0^{1/2}$. Moreover, $f_t \stackrel{i.i.d.}{\sim} \mathcal{N}(0, 1)$ and $U = \epsilon A$, where ϵ is an $N \times T$ matrix with i.i.d. $\mathcal{N}(0, 1)$ entries and A is a $T \times T$ matrix satisfying $\|A\| \lesssim 1$.

3.4.1 Principal Component Regression

Formally, we present the algorithm below:

Algorithm 2 (PCA Regression).

Inputs: $\bar{Y}$, $\underline{X}$, $\underline{W}$, x_T , and w_T .

S1. Apply SVD on $\underline{X}\mathbb{M}_{\underline{W}'}$ and obtain the estimated factors $\hat{\underline{F}}_{PCA} = \hat{\zeta}'\underline{X}\mathbb{M}_{\underline{W}'}$, where $\hat{\zeta} \in \mathbb{R}^{N \times K}$ are the first K left singular vectors of $\underline{X}\mathbb{M}_{\underline{W}'}$. Estimate the coefficients $\hat{\alpha} = \bar{Y}\hat{\underline{F}}_{PCA}' \left(\hat{\underline{F}}_{PCA}\hat{\underline{F}}_{PCA}' \right)^{-1}$.

S2. Obtain $\hat{\gamma} = \hat{\alpha}\hat{\zeta}'$ and output the prediction $\hat{y}_{T+h}^{PCA} = \hat{\gamma}x_T + (\hat{\alpha}_w - \hat{\gamma}\hat{\beta}_w)w_T$, where $\hat{\alpha}_w = \bar{Y}\underline{W}'(\underline{W}\underline{W}')^{-1}$ and $\hat{\beta}_w = \underline{X}\underline{W}'(\underline{W}\underline{W}')^{-1}$.

Outputs: $\hat{y}_{T+h}^{PCA}$, $\hat{\underline{F}}_{PCA}$, $\hat{\alpha}$, $\hat{\alpha}_w$, $\hat{\beta}_w$, and $\hat{\gamma}$.

Proposition 1. In Example 4, suppose that $N/(N_0T) \rightarrow \delta \geq 0$ and $\|\beta\| \rightarrow \infty$ and define M as $M := T^{-1}\underline{F}'\underline{F} + \delta A_1'A_1$, where A_1 is the first $T-h$ columns of A . Then, if the two leading eigenvalues of M are distinct in the sense that $(\lambda_1(M) - \lambda_2(M))/\lambda_1(M) \gtrsim_P 1$, the estimated factor $\hat{\underline{F}}_{PCA}$ satisfies

$$\left\| \mathbb{P}_{\hat{\underline{F}}_{PCA}'} - \mathbb{P}_{\eta_{PCA}} \right\| \xrightarrow{P} 0,$$

where η_{PCA} is the first eigenvector of M . In the special case that $A_1' A_1 = \mathbb{I}_{T-h}$, it satisfies that

$$\left\| \mathbb{P}_{\hat{\underline{F}}_{PCA}'} - \mathbb{P}_{\underline{F}'} \right\| \xrightarrow{P} 0.$$

Proposition 1 first shows that even if the number of factors is known to be 1, the factor estimated by PCA is in general inconsistent, because the eigenvector η_{PCA} deviates from that of $T^{-1} \underline{F}' \underline{F}$, as the latter is polluted by A . In the special case where error is homoskedastic and has no serial correlation, i.e., $A_1' A_1 = \mathbb{I}_{T-h}$, the estimated factor becomes consistent, in that $\delta A_1' A_1$ in M does not change the eigenvectors of $T^{-1} \underline{F}' \underline{F}$. This result echoes a similar result in Section 4 of Bai (2003), who established the consistency of factors with homoskedasticity and serially independent error even when T is fixed. That said, while factors can be estimated consistently in this special case, the prediction of y_{T+h} based on Algorithm 2 is not consistent.

Proposition 2. *Under the same assumptions as in Proposition 1, if we further assume $A_1' A_1 = \mathbb{I}_{T-h}$, then we have $\hat{y}_{T+h}^{PCA} \xrightarrow{P} (1 + \delta)^{-1} E_T(y_{T+h})$.*

The reason behind the inconsistency is that even though $\hat{\underline{F}}_{PCA}$, (effectively the right singular vector of $\underline{X}$) is consistent in the special case, the left singular vector, $\hat{\varsigma}$ and the singular values are not consistent, which lead to a biased prediction. This result demonstrates the limitation of PC regressions in the presence of weak factor structure.

3.4.2 Partial Least Squares

PCA is an unsupervised approach, in that the PCs are obtained without any information from the prediction target. Therefore, it might be misled by large idiosyncratic errors in x_t when the signal is not sufficiently strong. In contrast with PCA, partial least squares (PLS) is another supervised technique for prediction, which has been shown to work better than PCA in other settings, see, e.g., Kelly and Pruitt (2013). Unlike PCA, PLS uses the information of the response variable when estimating factors. Ahn and Bae (2022) develop its asymptotic properties for prediction in the case of strong factors. We now investigate its asymptotic performance in the same setting above.

The PLS regression algorithm is formulated below:

Algorithm 3 (PLS). *The estimator proceeds as follows:*

Inputs: $\bar{Y}$, $\underline{X}$, $\underline{W}$, x_T , and w_T . *Initialization:* $Y_{(1)} := \bar{Y} \mathbb{M}_{\underline{W}'}$, $X_{(1)} := \underline{X} \mathbb{M}_{\underline{W}'}$.

S1. For $k = 1, 2, \dots, K$, repeat the following steps using $X_{(k)}$.

- a. Obtain the weight vector $\hat{\varsigma}_{(k)}$ from the largest left singular vector of $X_{(k)} Y_{(k)}'$.*
- b. Estimate the k th factor as $\hat{\underline{F}}_{(k)} = \hat{\varsigma}_{(k)}' X_{(k)}$.*
- c. Estimate coefficients $\hat{\alpha}_{(k)} = Y_{(k)} \hat{\underline{F}}_{(k)}' \left(\hat{\underline{F}}_{(k)} \hat{\underline{F}}_{(k)}' \right)^{-1}$ and $\hat{\beta}_{(k)} = X_{(k)} \hat{\underline{F}}_{(k)}' \left(\hat{\underline{F}}_{(k)} \hat{\underline{F}}_{(k)}' \right)^{-1}$.*

e. Remove $\hat{\underline{F}}_{(k)}$ to obtain residuals for the next step: $X_{(k+1)} = X_{(k)} - \hat{\beta}_{(k)}\hat{\underline{F}}_{(k)}$ and $Y_{(k+1)} = Y_{(k)} - \hat{\alpha}_{(k)}\hat{\underline{F}}_{(k)}$.

S2. Obtain $\hat{\gamma} = \hat{\alpha}\hat{\zeta}'$ and the prediction $\hat{y}_{T+h}^{PLS} = \hat{\gamma}x_T + (\hat{\alpha}_w - \hat{\gamma}\hat{\beta}_w)w_T$, where $\hat{\alpha}_w = \bar{Y}\underline{W}'(\underline{W}\underline{W}')^{-1}$ and $\hat{\beta}_w = \underline{X}\underline{W}'(\underline{W}\underline{W}')^{-1}$.

Outputs: $\hat{y}_{T+h}^{PLS}$, $\hat{\underline{F}}_{PLS} := (\hat{\underline{F}}'_{(1)}, \dots, \hat{\underline{F}}'_{(\hat{K})})'$, $\hat{\alpha}$, $\hat{\alpha}_w$, $\hat{\beta}_w$, and $\hat{\gamma}$.

The PLS estimator has a closed-form formula if Y is a $1 \times T$ vector and a single factor model is estimated ($K = 1$):

$$\hat{y}_{T+h}^{PLS} = \|\bar{Y}\underline{X}'\underline{X}\|^{-2}\bar{Y}\underline{X}'\underline{X}\bar{Y}'\bar{Y}\underline{X}'x_T.$$

While the PLS procedure is intuitively appealing, the next propositions show that this approach produces biased prediction results in the presence of weak factors.

Proposition 3. *In Example 4, suppose that $N/(N_0T) \rightarrow \delta \geq 0$ and $\|\beta\| \rightarrow \infty$, then the estimated factor $\hat{\underline{F}}_{PLS}$ satisfies*

$$\left\| \mathbb{P}_{\hat{\underline{F}}_{PLS}} - \mathbb{P}_{\eta_{PLS}} \right\| \xrightarrow{P} 0,$$

where $\eta_{PLS} = (\mathbb{I}_{T-h} + \delta A_1' A_1) \underline{F}'$. In the special case that $A_1' A_1 = \mathbb{I}_{T-h}$, it satisfies

$$\left\| \mathbb{P}_{\hat{\underline{F}}_{PLS}} - \mathbb{P}_{\underline{F}'} \right\| \xrightarrow{P} 0.$$

Proposition 4. *Under the assumptions of Proposition 3, if we further assume that $A_1' A_1 = \mathbb{I}_{T-h}$, then we have $\hat{y}_{T+h}^{PLS} \xrightarrow{P} (1 + \delta)^{-1} E_T(y_{T+h})$.*

Therefore, the consistency of the PLS factor also depends on the homoskedasticity assumption $A_1' A_1 = \mathbb{I}_{T-h}$ and the forecasting performance of PLS regression is similar to PCA in our weak factor setting. The reason is that the information about the covariance between $\underline{X}$ and $\bar{Y}$ used by PLS is dominated by the noise component of $\underline{X}$, hence PLS does not resolve the issue of weak factors, despite it being a supervised predictor.

Finally, before we conclude the analysis on PLS, we demonstrate a potential issue of PLS due to “overfitting.” It turns out that PLS can severely overfit the in-sample data and perform badly out of sample, because PLS overuses information on y to construct its predictor. We illustrate this issue with the following example:

EXAMPLE 5: Suppose x_t and y_{t+h} follow a “0-factor” model:

$$x_t = u_t, \quad y_{t+h} = z_{t+h},$$

where u_t s follow i.i.d. $\mathcal{N}(0, \mathbb{I}_N)$ and z_t s follow i.i.d. $\mathcal{N}(0, 1)$.

Proposition 5. *In Example 5, if we use $\hat{K} = 1$, then we have $\hat{y}_{T+h}^{PLS} \gtrsim_P N^{3/2}T^{1/2}/(N^2 + T^2)$ while $\hat{y}_{T+h}^{PCA} \lesssim_P 1/(N^{1/2} + T^{1/2})$. Specifically, in the case of $N \asymp T$, $\hat{y}_{T+h}^{PLS} \gtrsim_P 1$ and $\hat{y}_{T+h}^{PCA} \lesssim_P N^{-1/2}$.*

The conditional expectation of y_{T+h} is 0 in this example, but $\hat{y}_{T+h}^{PLS}$ can be bounded away from 0 when using more factors than necessary. In contrast, $\hat{y}_{T+h}^{PCA}$ remains consistent. The failure of PLS is precisely due to that it selects a component in x that appears correlated with y , despite the fact that there is no correlation between them in this DGP. While SPCA's behavior is difficult to pin down in this example, intuitively, it falls in between these two cases. When q is very large, SPCA resembles PCA as it uses a large number of predictors in x to obtain components. When q is too small, SPCA is prone to overfitting like PLS. With a good choice of q by cross-validation, SPCA can also avoid overfitting.

3.4.3 PCA Regression of Stock and Watson (2002)

Stock and Watson (2002) adopt an alternative version of the PCA regression algorithm (hereafter SW-PCA) to what we have presented in Algorithm 2. The key difference is that SW-PCA conducts PCA on the entire X instead of $\underline{X}$. Therefore, they can obtain $\hat{f}_T$ directly from this step, instead of reconstructing it using the estimated weights in-sample. While our focus is not on PCA, the PCA algorithm is part of our SPCA procedure. Given the popularity of SW-PCA, we explain why we prefer our version of PCA regression given by Algorithm 2.

Formally, we present their algorithm below:

Algorithm 4 (SW-PCA).

Inputs: $\bar{Y}$, X , and W .

S1. Apply SVD on X , and obtain the estimated factors $\hat{F}_{SW} = \hat{\zeta}_* X \mathbb{M}_{W'}$, where $\hat{\zeta}_* \in \mathbb{R}^{N \times K}$ are the first K left singular vectors of X . Estimate the coefficients by time-series regression: $\hat{\alpha} = \bar{Y} \mathbb{M}_{\underline{W}'} \hat{F}_{SW}' \left(\hat{F}_{SW} \mathbb{M}_{\underline{W}'} \hat{F}_{SW}' \right)^{-1}$ ⁵ and $\hat{\alpha}_w = \bar{Y} \mathbb{M}_{\hat{F}_{SW}'} \underline{W}' \left(\underline{W} \mathbb{M}_{\hat{F}_{SW}'} \underline{W}' \right)^{-1}$.

S2. Obtain the prediction $\hat{y}_{T+h}^{SW} = \hat{\alpha} \hat{f}_T + \hat{\alpha}_w w_T$, where $\hat{f}_T$ is the last column of $\hat{F}_{SW}$ and $\hat{\alpha}_w = \bar{Y} \underline{W}' (\underline{W} \underline{W}')^{-1}$.

Outputs: $\hat{y}_{T+h}^{SW}$, $\hat{F}_{SW}$, $\hat{\alpha}$, and $\hat{\alpha}_w$.

The advantage of SW-PCA is that the consistency of factors is sufficient for the consistency of the prediction, unlike PCA as shown by Proposition 2. In other words, even though this is not true in general, $\hat{y}_{T+h}^{SW}$ can be consistent in the special case $A'A = \mathbb{I}_T$. Additionally, SW-PCA is more efficient for factor estimation in that it uses the entire data matrices X and W .

Nevertheless, the negative side of the SW-PCA is that it can be unstable because it is more prone to overfitting. We illustrated this issue using the example below.

⁵Unlike Algorithm 1, $\hat{F}_{SW}$ is not orthogonal to $\underline{W}$.

EXAMPLE 6: Suppose x_t and y_{t+h} follow a “0-factor” model:

$$x_t = u_t, \quad y_{t+h} = z_{t+h},$$

where u_t s are generated from mean zero normal distributions independently with $\text{Cov}(u_t) = \mathbb{I}_N$ for $t < T$ and $\text{Var}(u_T) = (1 + \epsilon)\mathbb{I}_N$ for some constant $\epsilon > 0$, and z_t s follow i.i.d. $\mathcal{N}(0, 1)$.

Proposition 6. *In Example 6, suppose that $T/N \rightarrow 0$, if we use $\hat{K} = 1$, then we have $\text{Var}(\hat{y}_{T+h}^{SW}) \rightarrow \infty$ and $\hat{y}_{T+h}^{PCA} \xrightarrow{P} 0$.*

Intuitively, SW-PCA uses in-sample estimates of the eigenvectors based on data up to T as factors for prediction, whereas PCA uses out-of-sample estimates of the factors, constructed at time T but based on weights estimated up to $T - h$. Because of this, SW-PCA may suffer more from “overfitting” compared to PCA, if the statistical properties of the data differ from $T - h$ to T . Example 6 investigates the case with heteroskedastic u_T in the scenario of overfitting $\hat{K} = 1 > K = 0$, in which case SW-PCA could perform rather wildly. This example appears contrived, but in practice macroeconomic data are often heterogenous and the number of factors is difficult to pin down. Such an issue is thereby relevant and we hence advocate Algorithms 2 for robustness.

3.5 Tuning Parameter Selection

Along with the gain in robustness to weak factors comes the cost of an extra tuning parameter. To implement the SPCA estimator, we need to select two tuning parameters, q and c . The parameter q dictates the size of the subset used for PCA construction, whereas the parameter c determines the stopping rule, and in turn the number of factors, K . By comparison, PCA and PLS, effectively, only require selecting K . We have established in Theorem 3 that we can consistently recover K , provided q and c satisfy certain conditions.

In practice, we may as well directly tune K instead of c , given that K is more interpretable, that K can only take integer values, and that the scree plot is informative about reasonable ranges of K . With respect to q , a larger choice of q renders the performance of SPCA resemble that of PCA, and hence becomes less robust to weak factors. Smaller values of q elevate the risk of overfitting, because the selected predictors are more prone to overfit y . We suggest tuning $\lfloor qN \rfloor$ instead of q , because the former can only take integer values, and that multiple choices of the latter may lead to the same integer values of the former.

In our applications, we select tuning parameters based on 3-fold cross-validation that proceeds as follows. We split the entire sample into 3 consecutive folds. Because of the time series dependence, we do not create these folds randomly. We then use each of the three folds, in turn, for validation while the other two are used for training. We select the optimal tuning parameters according to the average R^2 in the validation folds. With these selected parameters, we refit the model using the

entire data before making predictions. We conduct a thorough investigation of the effect of tuning on the finite sample performance of all procedures below.

4 Simulations

In this section, we study the finite sample performance of our SPCA procedure using Monte Carlo simulations.

Specifically, we consider a 3-factor DGP as given by equation (1) with two strong factors f_{1t} , f_{2t} and one potentially weak factor f_{3t} . For strong factors f_{1t} and f_{2t} , we generate exposure to them independently from $\mathcal{N}(0, 1)$. To simulate a weak factor f_{3t} , we generate exposure to it from a Gaussian mixture distribution, drawing values with probability a from $\mathcal{N}(0, 1)$ and $1 - a$ from $\mathcal{N}(0, 0.1^2)$. The parameter a determines the strength of the third factor and it ranges from $\{0.5, 0.1, 0.05\}$ in the simulations.

Our aim is to predict y_{T+1} , or equivalently, estimate $E_T(y_{T+1}) = \alpha f_T + \alpha_w w_T$, where w includes an intercept term and a lagged term of y . We consider two DGPs for y . In the first scenario, we set $\alpha_w = (0, 0.2)$ and $\alpha = (0, 0, 1)$, i.e., $y_{t+1} = f_{3t} + 0.2y_t + z_{t+1}$. Since y is a univariate target, there is no guarantee that we can recover all factors. We thus examine the consistency of the prediction, as shown in Theorem 2, on the basis of MSE and $\|\hat{\gamma}\beta - \alpha\|$. In the second scenario, we examine the quality of factor space recovery and inference. We thereby simulate a multivariate target with $\alpha = \mathbb{I}_3$ and $\alpha_w = (0_{3 \times 1}, 0.2\mathbb{I}_3)$, i.e., $y_{i,t+1} = f_{it} + 0.2y_{it} + z_{i,t+1}$, for $i = 1, 2, 3$.

We generate realizations of f_{it} , z_{it} independently from the standard normal distribution. To generate u_{it} , we first draw ϵ_{its} from $\mathcal{N}(0, 3)$ independently and construct the matrix $A = S\Gamma$, where S is a $(T + 1) \times (T + 1)$ diagonal matrix with elements drawn from $\text{Unif}(0.5, 1.5)$ and Γ is a $(T + 1) \times (T + 1)$ rotation matrix drawn uniformly from a unit sphere. Therefore, u_{it} as constructed by $U = \epsilon A$ features heteroskedasticity.

Table 1 compares the finite sample performance of SPCA, PCA, and PLS in the first scenario. In both panels, the sample size is $T = 60, 120$, and around $aN = 100$ predictors have exposure to the factor f_{3t} . We simulate $N = 200$ ($a = 0.5$) predictors in the upper panel, so that f_{3t} is exposed to half of them and is thereby strong, and set $N = 2,000$ ($a = 0.05$) in the lower panel, where f_{3t} becomes much weaker due to the large number of predictors that do not load on it.

To highlight the sensitivity of all estimators to the number of factors, we separately report results for each choice of K from 1 to 5 (not tuned), while only selecting the other tuning parameter q for SPCA via cross-validation. We also report results with both parameters tuned jointly for SPCA, and the single parameter K tuned for PCA and PLS, respectively.

The simulation results in Table 1 square well with our theoretical predictions. In the strong factor case (upper panel), PCA and SPCA perform similarly. They achieve minimum prediction error when K is set at the true value 3 in that the first two factors do not predict y . This suggests that tuning q does not worsen the performance of SPCA. PLS can also achieve desirable performance

but typically with K smaller than 3. Interestingly, its performance deteriorates rapidly as K increases and surpasses the true value. The reason, as we explain in Proposition 5, is that PLS is more likely to overfit as it uses information about y to directly construct predictors. In contrast, PCA based approaches are more robust to noisy factors used in prediction.

As to the weak factor case (lower panel), SPCA outperforms both PLS and PCA as predicted by our theory. Moreover, SPCA tends to achieve optimal performance when $K = 2$. Recall that in this case, we do not have asymptotic guarantee that SPCA can recover the entire factor space. For this reason, it is possible that a third factor out of this procedure contributes more noise than signal, hence the performance of SPCA deteriorates with an additional factor.

Both panels show that tuning K in most cases slightly deteriorates the optimal prediction MSE and estimation error. That said, the resulting errors remain smaller than what the second best choice of K can achieve.

Table 1: Finite Sample Comparison of Predictors (Univariate y)

		MSE						$\ \hat{\gamma}\beta - \alpha\ $					
	K	1	2	3	4	5	$\hat{K}$	1	2	3	4	5	$\hat{K}$
T		Panel A: $N = 200 \quad a = 0.5$											
60	SPCA	0.91	0.52	0.15	0.17	0.17	0.16	0.92	0.59	0.24	0.25	0.25	0.25
	PCA	1.05	1.08	0.15	0.15	0.15	0.15	1.01	1.02	0.26	0.26	0.25	0.26
	PLS	0.34	0.17	0.37	0.51	0.70	0.21	0.50	0.22	0.28	0.27	0.27	0.25
120	SPCA	0.89	0.49	0.09	0.11	0.11	0.10	0.92	0.55	0.17	0.17	0.17	0.17
	PCA	1.04	1.06	0.09	0.09	0.09	0.09	1.00	1.01	0.17	0.17	0.17	0.17
	PLS	0.25	0.10	0.31	0.40	0.66	0.11	0.38	0.16	0.26	0.18	0.19	0.16
		Panel B: $N = 2000 \quad a = 0.05$											
60	SPCA	0.75	0.29	0.41	0.52	0.58	0.36	0.78	0.32	0.42	0.45	0.47	0.36
	PCA	1.11	1.14	0.69	0.67	0.65	0.67	1.01	1.03	0.75	0.74	0.73	0.74
	PLS	1.14	0.55	0.52	0.67	0.75	0.55	1.00	0.56	0.50	0.49	0.47	0.51
120	SPCA	0.55	0.13	0.18	0.26	0.27	0.16	0.65	0.19	0.28	0.34	0.35	0.22
	PCA	1.05	1.08	0.27	0.27	0.27	0.27	1.01	1.02	0.44	0.44	0.44	0.44
	PLS	0.94	0.24	0.26	0.45	0.55	0.23	0.92	0.34	0.30	0.32	0.30	0.29

Notes: We evaluate the performance of SPCA, PCA, and PLS in terms of prediction MSE and $\|\hat{\gamma}\beta - \alpha\|$. All numbers reported are based on averages over 1,000 Monte Carlo repetitions. We highlight the best values based on each criterion in bold.

Furthermore, Table 2 reports the performance of SPCA, PCA, and PLS for each entry of y in the multi-target scenario. In this case we only report results with parameters tuned. As discussed previously, we expect the recovery of all factors using SPCA, because to each factor, at least one entry of y_t has exposure. We first report the distance between $\hat{F}$ and the true factors F , defined by $d(\hat{F}, F) = \|\mathbb{P}_{\hat{F}} - \mathbb{P}_F\|$. We also report the MSE _{i} s for $\hat{y}_{i,T+1}$, $i = 1, 2, 3$, where MSE₃ is based on $y_{3,T+1}$, which depends on the potentially weak factor f_{3T} by construction. Again, we vary the value of a and N , while maintaining $aN = 100$, so that the number of predictors with exposure to the third factor is fixed throughout.

The findings here are again consistent with our theory. In particular, as a varies from 0.5 to 0.05, the third factor becomes increasingly difficult to detect. Both PCA and PLS report a substantially larger distance $d(\hat{F}, F)$ than SPCA. In the mean-time, the distortion in the factor space translates to larger prediction errors for the third target y_3 , in that it loads on the weak factor f_3 besides its

own lag. Throughout this experiment, SPCA maintains almost the same level of performance as a varies, demonstrating its robustness to weak factors.

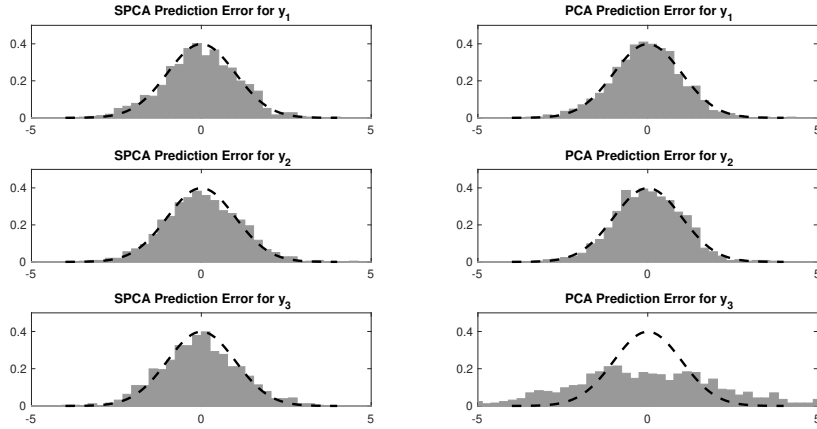
Table 2: Finite Sample Comparison of Predictors (Multivariate y)

a	SPCA				PCA				PLS			
	$d(\hat{F}, F)$	MSE ₁	MSE ₂	MSE ₃	$d(\hat{F}, F)$	MSE ₁	MSE ₂	MSE ₃	$d(\hat{F}, F)$	MSE ₁	MSE ₂	MSE ₃
$T = 60$												
0.5	0.40	0.14	0.16	0.20	0.40	0.14	0.15	0.21	0.41	0.14	0.15	0.19
0.1	0.44	0.13	0.14	0.25	0.55	0.12	0.12	0.55	0.54	0.12	0.13	0.38
0.05	0.45	0.14	0.13	0.27	0.66	0.12	0.11	0.72	0.59	0.12	0.12	0.53
$T = 120$												
0.5	0.30	0.07	0.08	0.10	0.30	0.07	0.08	0.10	0.31	0.08	0.08	0.10
0.1	0.31	0.07	0.07	0.12	0.36	0.06	0.06	0.22	0.39	0.07	0.06	0.17
0.05	0.31	0.07	0.07	0.11	0.39	0.06	0.06	0.29	0.44	0.06	0.06	0.22

Notes: We evaluate the performance of SPCA, PCA, and PLS in terms of the distance between estimated factor space and the true factor space, $d(\hat{F}, F) = \|\mathbb{P}_{\hat{F}'} - \mathbb{P}_{F'}\|$, as well as MSE _{i} for predicting the i th entry of y . All numbers reported are based on averages over 1,000 Monte Carlo repetitions. We vary the value a takes, while fixing $aN = 100$.

Last but not least, we report the histograms of the standardized prediction errors using the CLT of Theorem 4 in Figure 1. The setting is identical to that of Table 2 with $a = 0.05$ and $T = 120$. The histograms match well with the standard normal density for SPCA, and hence verifies the central limit result we derive. As to PCA, there is visible distortion to normality for y_3 , due to the presence of the weak factor f_3 .

Figure 1: Histograms of the Standardized Prediction Errors



Notes: We provide histograms of standardized prediction errors for each entry of y using SPCA and PCA, respectively, based on 1,000 Monte Carlo repetitions. The dashed curve on each plot corresponds to the standard normal density.

5 Conclusion

The problem of macroeconomic forecasting is central in both academic research as well as for designing policy. The availability of large datasets has spurred the development of methods, pioneered by

Stock and Watson (2002), aimed at reducing the dimensionality of the predictors in order to preserve parsimony and achieve better out of sample predictions.

The existing methods that are typically applied to this problem aim to extract a common predictive signal from the large set of available predictors, separating it from the noise and reducing the problem’s dimensionality. What our paper adds to this literature is the idea that the availability of a large number of predictors also allows us to *discard* predictors that are not sufficiently informative. That is, predictors that are mostly noise actually hurt the signal extraction because they contaminate the estimation of the common component contained in other, more informative, signals.

How can one know which predictors are noisy and which are useful? The key idea of SPCA is that one can discriminate between useful and noisy predictors by having the target itself guide the selection. This idea, first proposed in Bair and Tibshirani (2004), naturally leads to adding a screening step before factor extraction. But this original version of SPCA only works in very constrained environments that they can all be extracted via PCA from the same subset of predictors.

In practice, there is no guarantee for that to be the case. Whether a latent factor is strong or weak (and *how* strong) depends on how exposed the various predictors are to it – and each empirical applications could feature a different mix of strong and weak latent factors. Therefore, we propose a new SPCA approach that iterates a selection step, a factor extraction step, and a projection step. As we demonstrate in the paper, this procedure can consistently handle a whole range of latent factor strength. Our empirical analysis shows that indeed this procedure fares well in an application with a large number of potentially noisy macroeconomic predictors.

Two final points are worth noting. First, like any procedure, it will work best under some DGPs, and worse under others. In particular, the procedure will potentially miss factors that are extremely weak – no procedure can ever distinguish them from noise, because the exposures of the predictors to these factors are simply too small.

Second, our theory highlights an interesting tradeoff that emerges when working with weak factors. Detecting the weak factors using unsupervised methods (like PCA) is, by definition, difficult or impossible: there is a wide range of strength of factors that will be missed by these methods. Methods based on supervised selection can help extract additional signal, thanks to the guidance from the target. This ability comes at a cost: the possibility of missing factors that are not related to the target. Therefore, this procedure is most useful in applications, like forecasting, where omitting factors not related to the target does not bias the prediction. We leave to future work an additional exploration of other contexts in which SPCA can be useful.

A Mathematical Proofs

For notation simplicity, we use X , F , U , Y , Z in place of $\underline{X}$, $\underline{F}$, $\underline{U}$, $\overline{Y}$, and $\overline{Z}$, and use T_h for $T - h$. In addition, without loss of generality, we assume that $\Sigma_f = \mathbb{I}_K$ in the proof, in that we can always normalize the factors by $\Sigma_f^{-1/2}$ and redefine β in (1) and α in (2) accordingly.

A.1 Proof of Theorem 1

Proof. We start with the DGP without w_t first. Throughout the proof, we use $\tilde{X}_{(k)} := (X_{(k)})_{[\hat{I}_k]}$ to denote the matrix on which we perform SVD in each step of Algorithm 1. The first left and right singular vectors of $\tilde{X}_{(k)}$ are denoted by $\hat{\varsigma}_{(k)}$ and $\hat{\xi}_{(k)}$, while the largest singular value of $\tilde{X}_{(k)}$ is denoted by $\sqrt{T_h \hat{\lambda}_{(k)}}$. As a result, $\hat{\lambda}_{(k)} = T_h^{-1} \left\| \tilde{X}_{(k)} \right\|^2$. Moreover, by definition

$$\hat{\varsigma}_{(k)} = T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \tilde{X}_{(k)} \hat{\xi}_{(k)}, \quad \hat{\xi}_{(k)} = T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \tilde{X}_{(k)}' \hat{\varsigma}_{(k)}. \quad (12)$$

Therefore, our estimated factor at k -th step is $\hat{F}_{(k)} = \hat{\varsigma}_{(k)}' \tilde{X}_{(k)} = T_h^{1/2} \hat{\lambda}_{(k)}^{1/2} \hat{\xi}_{(k)}'$. Consequently, the coefficients of regressing X and Y onto this factor are, respectively:

$$\hat{\beta}_{(k)} = T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} X_{(k)} \hat{\xi}_{(k)} \quad \text{and} \quad \hat{\alpha}_{(k)} = T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} Y_{(k)} \hat{\xi}_{(k)}. \quad (13)$$

Then we define $\tilde{D}_{(k)} \in \mathbb{R}^{q \times N}$ iteratively by

$$\tilde{D}_{(k)} = (\mathbb{I}_N)_{[\hat{I}_k]} - \sum_{i=1}^{k-1} T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} X_{[\hat{I}_k]} \hat{\xi}_{(i)} \hat{\varsigma}_{(i)}' \tilde{D}_{(i)},$$

with $\tilde{D}_{(1)} = (\mathbb{I}_N)_{[\hat{I}_1]}$. We can show by induction that $\tilde{X}_{(k)} = \tilde{D}_{(k)} X$. In fact, by Lemma 1, we have $\hat{\xi}_{(i)}' \hat{\xi}_{(j)} = 0$ for $i \neq j \leq \hat{K}$ which suggests that $\hat{F}_{(k)}$'s for all k are pairwise orthogonal. Using this property and the definition of $\tilde{X}_{(k)}$, we have

$$\tilde{X}_{(k)} = (X_{(k)})_{[\hat{I}_k]} = X_{[\hat{I}_k]} \prod_{i=1}^{k-1} \mathbb{M}_{\hat{F}_{(i)}} = X_{[\hat{I}_k]} \left(\mathbb{I}_{T_h} - \sum_{i=1}^{k-1} \hat{\xi}_{(i)} \hat{\xi}_{(i)}' \right), \quad (14)$$

for $k > 1$ and when $k = 1$,

$$\tilde{X}_{(1)} = X_{[\hat{I}_1]} = \beta_{[\hat{I}_1]} F + U_{[\hat{I}_1]}.$$

Using (12), if $\tilde{X}_{(i)} = \tilde{D}_{(i)} X$ for any $i < k$ we can write (14) as

$$\tilde{X}_{(k)} = X_{[\hat{I}_k]} \left(\mathbb{I}_{T_h} - \sum_{i=1}^{k-1} \hat{\xi}_{(i)} \hat{\xi}_{(i)}' \right) = X_{[\hat{I}_k]} - \sum_{i=1}^{k-1} T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} X_{[\hat{I}_k]} \hat{\xi}_{(i)} \hat{\varsigma}_{(i)}' \tilde{X}_{(i)} = \tilde{D}_{(k)} X.$$

Since $\tilde{X}_{(1)} = X_{[\hat{I}_1]} = \tilde{D}_{(1)} X$ holds immediately by definition, we have $\tilde{X}_{(k)} = \tilde{D}_{(k)} X$ by induction. In light of this, the estimated factors satisfy

$$\hat{F}_{(k)} = \hat{\varsigma}_{(k)}' \tilde{X}_{(k)} = \hat{\varsigma}_{(k)}' \tilde{D}_{(k)} X, \quad (15)$$

for all k , and by definition, we have $\hat{\zeta}_{(k)} = (\hat{\zeta}'_{(k)} \tilde{D}_{(k)})'$. Moreover, using (13) the estimated coefficient $\hat{\gamma}$ can be written as

$$\hat{\gamma} = \sum_{k=1}^{\hat{K}} \hat{\alpha}_{(k)} \hat{\zeta}'_{(k)} \tilde{D}_{(k)} = \sum_{k=1}^{\hat{K}} T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} Y \hat{\xi}_{(k)} \hat{\zeta}'_{(k)} \tilde{D}_{(k)}. \quad (16)$$

We further define $\tilde{\beta}_{(k)} = \tilde{D}_{(k)} \beta$ and $\tilde{U}_{(k)} = \tilde{D}_{(k)} U$, then $\tilde{X}_{(k)}$ can be written in the form of

$$\tilde{X}_{(k)} = \tilde{\beta}_{(k)} F + \tilde{U}_{(k)}. \quad (17)$$

We also define the population analog of $\tilde{D}_{(k)}$ for each k by

$$D_{(k)} = (\mathbb{I}_N)_{[I_k]} - \sum_{i=1}^{k-1} \lambda_{(i)}^{-1/2} \beta_{[I_k]} b_{(i)} \varsigma'_{(i)} D_{(i)}, \quad D_{(1)} = (\mathbb{I}_N)_{[I_1]},$$

where $\sqrt{\lambda_{(k)}}$ is the leading singular value of $\beta_{(k)}$, $\varsigma_{(k)}$ and $b_{(k)}$ are the corresponding left and right singular vectors of $\beta_{(k)}$. By a similar induction argument, we can show that

$$\beta_{(k)} = \beta_{[I_k]} \prod_{i < k} \mathbb{M}_{b_{(i)}} = D_{(k)} \beta.$$

Intuitively, $\tilde{\beta}_{(k)}$ and $\tilde{D}_{(k)}$ are sample analogs of $\beta_{(k)}$ and $D_{(k)}$.

Similar representations to (17) can be constructed for $Y_{(k)} := Y \prod_{i=1}^{k-1} \mathbb{M}_{\hat{F}'_{(i)}}$ for each k . Specifically, we have

$$Y_{(k)} = Y \left(\mathbb{I}_{T_h} - \sum_{i=1}^{k-1} \hat{\xi}_{(i)} \hat{\xi}'_{(i)} \right) = \tilde{\alpha}_{(k)} F + \tilde{Z}_{(k)}, \quad (18)$$

where $\tilde{\alpha}_{(k)} \in \mathbb{R}^{D \times K}$ and $\tilde{\beta}_{(k)} \in \mathbb{R}^{|\hat{I}_k| \times K}$ are defined as

$$\tilde{\alpha}_{(k)} := \alpha - \sum_{i=1}^{k-1} T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} Y \hat{\xi}_{(i)} \hat{\zeta}'_{(i)} \tilde{\beta}_{(i)} \text{ and } \tilde{Z}_{(k)} := Z - \sum_{i=1}^{k-1} T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} Y \hat{\xi}_{(i)} \hat{\zeta}'_{(i)} \tilde{U}_{(i)}.$$

By Lemma 3, we have $P(\hat{I}_k = I_k) \rightarrow 1$ for $k \leq \tilde{K}$ and $P(\hat{K} = \tilde{K}) \rightarrow 1$. Thus, with probability approaching one, we can impose that $\hat{I}_k = I_k$ for any k and $\hat{K} = \tilde{K}$ in what follows.

To prove Theorem 1, using (17), the estimated factors can be written as

$$\hat{F}_{(k)} = \hat{\zeta}'_{(k)} \tilde{X}_{(k)} = \hat{\zeta}'_{(k)} \tilde{\beta}_{(k)} F + \hat{\zeta}'_{(k)} \tilde{U}_{(k)}.$$

Using Lemma 5(i), $\|\widehat{F}_{(k)}\| = \sqrt{T_h \widehat{\lambda}_{(k)}}$, and $\|\mathbb{M}_{F'}\| \leq 1$, we have

$$\|\widehat{F}_{(k)}\|^{-1} \|\widehat{F}_{(k)} \mathbb{M}_{F'}\| \leq \|\widehat{F}_{(k)}\|^{-1} \|\zeta'_{(k)} \widetilde{U}_{(k)}\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}.$$

□

A.2 Proof of Theorem 2

Proof. By definition of $X_{(k)}$ in Algorithm 1, we have

$$X_{(k)} = X_{(k-1)} \mathbb{M}_{\widehat{F}'_{(k-1)}} = X \prod_{i=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(i)}} = X \left(\mathbb{I}_{T_h} - \sum_{i=1}^{k-1} \widehat{\xi}_{(i)} \widehat{\xi}'_{(i)} \right).$$

Therefore, using (18), we have

$$X_{(k)} Y'_{(k)} = X \left(\mathbb{I}_{T_h} - \sum_{i=1}^{k-1} \widehat{\xi}_{(i)} \widehat{\xi}'_{(i)} \right) Y'_{(k)} = X Y'_{(k)}$$

as $Y_{(k)} \widehat{\xi}_{(i)} = 0$ for $i < k$ by Lemma 1. Therefore, the covariance $(X_{(k)})_{[i]} Y'_{(k)}$ for each predictor equals to $X_{[i]} Y'_{(k)}$. Based on the stopping rule, if our algorithm stops at $\tilde{K}$, there are at most $qN - 1$ predictors among all satisfying $T_h^{-1} \|X_{[i]} Y'_{(\tilde{K}+1)}\|_{\text{MAX}} \geq c$. Let S denote the set of these predictors. For $i \in S^c$, we have

$$\|T_h^{-1} X_{[i]} Y'_{(\tilde{K}+1)}\|_{\text{F}}^2 \lesssim \|T_h^{-1} X Y'_{(\tilde{K}+1)}\|_{\text{MAX}}^2 \lesssim_P 1, \quad (19)$$

where we use $\|\beta\|_{\text{MAX}} \lesssim 1$ from Assumption 2 and Lemma 3(vi) in the last step. On the other hand, in light of the set I_0 in Assumption 2, we have

$$\begin{aligned} \sum_{i \in I_0} \|T_h^{-1} X_{[i]} Y'_{(\tilde{K}+1)}\|_{\text{F}}^2 &= \sum_{i \in I_0 \cap S} \|T_h^{-1} X_{[i]} Y'_{(\tilde{K}+1)}\|_{\text{F}}^2 + \sum_{i \in I_0 \cap S^c} \|T_h^{-1} X_{[i]} Y'_{(\tilde{K}+1)}\|_{\text{F}}^2 \\ &\lesssim_P |I_0 \cap S| + |I_0 \cap S^c| c^2 \leq qN + c^2 N_0 = o(N_0), \end{aligned} \quad (20)$$

where we use (19), $|S| \leq qN - 1$, $c \rightarrow 0$, and $qN/N_0 \rightarrow 0$. Consequently, (20) leads to $\|Y_{(\tilde{K}+1)} X'_{[I_0]}\| = o_P(N_0)$. Moreover, using (18) and that $X = \beta F + U$, we can decompose

$$Y_{(\tilde{K}+1)} X'_{[I_0]} = \tilde{\alpha}_{(\tilde{K}+1)} F F' \beta'_{[I_0]} + \tilde{\alpha}_{(\tilde{K}+1)} F U'_{[I_0]} + \tilde{Z}_{(\tilde{K}+1)} F' \beta'_{[I_0]} + \tilde{Z}_{(\tilde{K}+1)} U'_{[I_0]}. \quad (21)$$

Using (20), (21), Lemma 9(i)(ii), and the fact that $\|\beta_{[I_0]}\| \lesssim N_0^{1/2}$, we have

$$\|\tilde{\alpha}_{(\tilde{K}+1)} (F F' \beta'_{[I_0]} + F U'_{[I_0]})\| = o_P(N_0^{1/2} T). \quad (22)$$

Also, using Assumption 4(i), Assumption 1(i) and Weyl's theorem, we have

$$\begin{aligned} |\sigma_K(FF'\beta'_{[I_0]} + FU'_{[I_0]}) - \sigma_K(T_h\beta_{[I_0]})| &\leq \|FU'_{[I_0]}\| + \|T_h^{-1}FF' - \mathbb{I}_K\| \|T_h\beta_{[I_0]}\| \\ &\lesssim_P N_0^{1/2} T^{1/2}. \end{aligned} \quad (23)$$

Since Assumption 2 implies that $\sigma_K(\beta_{[I_0]}) \asymp N_0^{1/2}$, we have $\sigma_K(FF'\beta'_{[I_0]} + FU'_{[I_0]}) \asymp N_0^{1/2} T$. Using this result, (22) and the inequality $\|\tilde{\alpha}_{(\tilde{K}+1)}(FF'\beta'_{[I_0]} + FU'_{[I_0]})\| \geq \sigma_K(FF'\beta_{[I_0]} + FU'_{[I_0]}) \|\tilde{\alpha}_{(\tilde{K}+1)}\|$, we have $\|\tilde{\alpha}_{(\tilde{K}+1)}\| \xrightarrow{P} 0$. That is, by definition of $\tilde{\alpha}_{(\tilde{K}+1)}$ in (18),

$$\left\| \alpha - \sum_{i=1}^{\tilde{K}} Y \hat{\xi}_{(i)} \frac{\zeta'_{(i)} \tilde{\beta}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| = o_P(1). \quad (24)$$

Next, (16) and $\tilde{\beta}_{(k)} = \tilde{D}_{(k)}\beta$ imply that

$$\hat{\gamma}\beta = \sum_{i=1}^{\tilde{K}} T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} Y \hat{\xi}_{(i)} \zeta'_{(i)} \tilde{\beta}_{(i)}.$$

Therefore, (24) is equivalent to $\|\hat{\gamma}\beta - \alpha\| = o_P(1)$.

As shown in Lemma 12, Assumptions 1, 3, and 4 hold when we replace F , Z and U by $F\mathbb{M}_{W'}$, $Z\mathbb{M}_{W'}$ and $U\mathbb{M}_{W'}$. Therefore all of the lemmas and the result $\|\hat{\gamma}\beta - \alpha\| = o_P(1)$ also hold when w_t is included. We write the prediction error of y_{T+h} as

$$\begin{aligned} \hat{y}_{T+h} - \mathbb{E}_T(y_{T+h}) &= \hat{\gamma}x_T + (\hat{\alpha}_w - \hat{\gamma}\hat{\beta}_w)w_T - \alpha f_T - \alpha_w w_T \\ &= (\hat{\gamma}\beta - \alpha)(f_T - FW'(WW')^{-1}w_T) + \hat{\gamma}(u_T - UW'(WW')^{-1}w_T) + ZW'(WW')^{-1}w_T. \end{aligned} \quad (25)$$

Using (16) and $\|Y\| \leq \|\alpha F\| + \|Z\| \lesssim_P T^{1/2}$ by Assumption 1, we have

$$\|\hat{\gamma}u_T\| \leq \sum_{k \leq \tilde{K}} T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \|Y\| \|\hat{\xi}_{(k)}\| \|\zeta'_{(k)} \tilde{D}_{(k)} u_T\| \lesssim_P \sum_{k \leq \tilde{K}} \hat{\lambda}_{(k)}^{-1/2} \|\zeta'_{(k)} \tilde{D}_{(k)} u_T\|, \quad (26)$$

and

$$\begin{aligned} T_h^{-1} \|\hat{\gamma}UW'\| &\leq \sum_{k \leq \tilde{K}} T_h^{-3/2} \hat{\lambda}_{(k)}^{-1/2} \|Y\| \|\hat{\xi}_{(k)}\| \|\zeta'_{(k)} \tilde{D}_{(k)} UW'\| \\ &\lesssim_P \sum_{k \leq \tilde{K}} T_h^{-1} \hat{\lambda}_{(k)}^{-1/2} \|\zeta'_{(k)} \tilde{U}_{(k)} W'\|. \end{aligned} \quad (27)$$

Using $\widehat{\lambda}_{(k)} \asymp_P qN$ from Lemma 3 and Lemma 5(ii)(iv), we have

$$T_h^{-1} \widehat{\lambda}_{(k)}^{-1/2} \left\| \widehat{\varsigma}_{(k)} \widetilde{U}_{(k)} W' \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}, \quad \widehat{\lambda}_{(k)}^{-1/2} \left\| \widetilde{\varsigma}'_{(k)} \widetilde{D}_{(k)} u_T \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}. \quad (28)$$

Therefore, $\|\widehat{\gamma} u_T\| = o_P(1)$. Furthermore, with $\|(WW')^{-1}\| \lesssim_P T^{-1}$ from Assumption 1, we have $\|\widehat{\gamma} U W' (W W')^{-1}\| = o_P(1)$. Together with $\|F W'\| \lesssim_P T^{1/2}$, $\|Z W'\| \lesssim_P T^{1/2}$ from Assumption 1 and $\|\widehat{\gamma} \beta - \alpha\| = o_P(1)$, we show that each term of (25) vanishes, and hence $\widehat{y}_{T+h} - \mathbb{E}_T[y_{T+h}] \xrightarrow{P} 0$. $\square$

A.3 Proof of Theorem 3

Proof. As in the proof of Theorem 1, we impose that $\widehat{K} = \tilde{K}$ and $\widehat{I}_k = I_k$, since Lemma 3 shows that both events occur with probability approaching 1. As shown in Lemma 2(iv), under the assumption that $\lambda_K(\alpha' \alpha) \gtrsim 1$, we have $\tilde{K} = K$. Together with $P(\widehat{K} = \tilde{K}) \rightarrow 1$, we have obtained (i) of Theorem 3. Below we can directly impose that $\widehat{K} = K$.

Again, following the same argument above (25), we only need analyze the case without w_t . As $\widehat{F}_{(k)} = T_h^{1/2} \widehat{\lambda}_{(k)}^{1/2} \widehat{\xi}'_{(k)}$, Theorem 1 implies $\left\| \widehat{\xi}'_{(k)} \mathbb{M}_{F'} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}$ for $k \leq K$. Let v denote $F'(F F')^{-1/2}$, we have

$$\left\| \widehat{\xi} - \mathbb{P}_{F'} \widehat{\xi} \right\| = \left\| \widehat{\xi} - v v' \widehat{\xi} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}, \quad (29)$$

where $\widehat{\xi}$ is a $T \times K$ matrix with each column equal to $\widehat{\xi}_{(k)}$. (29) implies that $\left\| \widehat{\xi}' v v' \widehat{\xi} - \mathbb{I}_K \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}$. By Weyl's inequality, $|\sigma_i(\widehat{\xi}' v) - 1| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}$, for $1 \leq i \leq K$, and thus

$$\begin{aligned} \left\| v - \widehat{\xi} \widehat{\xi}' v \right\| &\leq \sigma_K^{-1}(v' \widehat{\xi}) \left\| v v' \widehat{\xi} - \widehat{\xi} \widehat{\xi}' v v' \widehat{\xi} \right\| \lesssim_P \left\| v v' \widehat{\xi} - \widehat{\xi} \right\| + \left\| \widehat{\xi} (\widehat{\xi}' v v' \widehat{\xi} - \mathbb{I}_K) \right\| \\ &\lesssim_P q^{-1/2} N^{-1/2} + T^{-1}. \end{aligned}$$

Then, using this, (29), and the fact that $\|v\| = 1$ and $\|\widehat{\xi}\| = 1$, we have

$$\left\| \mathbb{P}_{\widehat{F}'} - \mathbb{P}_{F'} \right\| = \left\| \widehat{\xi} \widehat{\xi}' - v v' \right\| \leq \left\| \widehat{\xi} (\widehat{\xi} - v v' \widehat{\xi})' \right\| + \left\| (\widehat{\xi} \widehat{\xi}' v - v) v' \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}.$$

Next, we need a more intricate analysis of $\widehat{\gamma}$. Recall from the proof of Theorem 2 that

$$\widehat{\gamma} \beta = \sum_{k=1}^{\tilde{K}} T_h^{-1/2} \widehat{\lambda}_{(k)}^{-1/2} Y \widehat{\xi}_{(k)} \widetilde{\varsigma}'_{(k)} \widetilde{\beta}_{(k)}. \quad (30)$$

Denote $B_1 = (b_{11}, \dots, b_{\widehat{K}1}) \in \mathbb{R}^{K \times \widehat{K}}$, $B_2 = (b_{12}, \dots, b_{\widehat{K}2}) \in \mathbb{R}^{K \times \widehat{K}}$, where

$$b_{k1} = T^{-1/2} F \widehat{\xi}_{(k)}, \quad b_{k2} = \widehat{\lambda}_{(k)}^{-1/2} \widetilde{\beta}'_{(k)} \widehat{\xi}_{(k)}. \quad (31)$$

By Lemma 6,

$$\left\| T_h^{-1/2} Z \widehat{\xi}_{(k)} - T_h^{-1} Z F' b_{k2} \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}. \quad (32)$$

As we impose that $\widehat{K} = \tilde{K} = K$, combining (30), (31) and (32), with $\|B_1\| \lesssim_P 1$, $\|B_2\| \lesssim_P 1$ from Lemma 10, we have

$$\left\| \widehat{\gamma} \beta - \alpha B_1 B_2' - T_h^{-1} Z F' B_2 B_2' \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}. \quad (33)$$

Using Lemma 10(iv)(v), we obtain $\left\| \widehat{\gamma} \beta - \alpha - T_h^{-1} Z F' \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}$. $\square$

A.4 Proof of Theorem 4

Proof. As in the proof of Theorem 2, we have $\left\| F W' (W W')^{-1} \right\| \lesssim_P T^{-1/2}$ from Assumption 1 and $\left\| \widehat{\gamma} U W' (W W')^{-1} \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}$ as shown in (27) and (28). Together with $\left\| \widehat{\gamma} \beta - \alpha - T_h^{-1} Z F' \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}$, we can derive from (25) that:

$$\widehat{y}_{T+h} - E_T(y_{T+h}) = T_h^{-1} Z F' f_T + Z W' (W W')^{-1} w_T + \widehat{\gamma} u_T + O_P(T^{-1} + q^{-1} N^{-1}).$$

By Assumption 1, we have $|\lambda_i \left(T_h^{-1} \Sigma_w^{-1/2} W W' \Sigma_w^{-1/2} \right) - 1| \lesssim_P T^{-1/2}$ and thus

$$\begin{aligned} & \left\| Z W' (W W')^{-1} w_T - T_h^{-1} Z W' \Sigma_w^{-1} w_T \right\| \leq T_h^{-1} \left\| Z W' \right\| \left\| (T_h^{-1} W W')^{-1} - \Sigma_w^{-1} \right\| \|w_T\| \\ & \lesssim_P T_h^{-1/2} \left\| T_h^{-1} \Sigma_w^{-1/2} W W' \Sigma_w^{-1/2} - \mathbb{I}_D \right\| = T_h^{-1/2} \max_{i \leq D} |\lambda_i \left(T_h^{-1} \Sigma_w^{-1/2} W W' \Sigma_w^{-1/2} \right) - 1| \\ & \lesssim_P T^{-1}. \end{aligned} \quad (34)$$

For $\widehat{\gamma} u_T$, by (16), we have $\widehat{\gamma} u_T = \sum_{k=1}^K \widehat{\alpha}_{(k)} \zeta'_{(k)} \widetilde{D}_{(k)} u_T$ and thus

$$\left\| \widehat{\gamma} u_T - \sum_{k=1}^K \lambda_{(k)}^{-1/2} \alpha b_{(k)} \zeta'_{(k)} D_{(k)} u_T \right\| \leq \sum_{k=1}^K \left\| \widehat{\alpha}_{(k)} \zeta'_{(k)} \widetilde{D}_{(k)} u_T - \lambda_{(k)}^{-1/2} \alpha b_{(k)} \zeta'_{(k)} D_{(k)} u_T \right\|. \quad (35)$$

Lemma 8(vi) gives

$$q^{-1/2} N^{-1/2} |\zeta'_{(k)} \widetilde{D}_{(k)} u_T - \zeta'_{(k)} D_{(k)} u_T| \lesssim_P T^{-1} + q^{-1} N^{-1}. \quad (36)$$

In addition, (13) and Lemma 1 give $\widehat{\lambda}_{(k)}^{1/2} \widehat{\alpha}_{(k)} = T_h^{-1/2} Y \widehat{\xi}_{(k)} = \alpha b_{k1} + T_h^{-1/2} Z \widehat{\xi}_{(k)}$. With (32), $\|Z F'\| \lesssim_P T^{1/2}$ and $\|b_{k2}\| \lesssim_P 1$ from Lemma 10(i), this equation leads to

$$\left\| \widehat{\lambda}_{(k)}^{1/2} \widehat{\alpha}_{(k)} - \alpha b_{k1} \right\| \leq \left\| T_h^{-1/2} Z \widehat{\xi}_{(k)} - T_h^{-1} Z F' b_{k2} \right\| + \left\| T_h^{-1} Z F' b_{k2} \right\| \lesssim_P T^{-1/2} + q^{-1} N^{-1}.$$

Using $\|b_{k2} - b_{(k)}\| \lesssim_P T^{-1/2} + q^{-1/2}N^{-1/2}$ implied by Lemma 10(iii) and $\widehat{\lambda}_{(k)} \asymp_P qN$ from Lemma 3(iii), we have

$$\begin{aligned} \left\| \widehat{\alpha}_{(k)} - \widehat{\lambda}_{(k)}^{-1/2} \alpha b_{(k)} \right\| &\leq \left\| \widehat{\alpha}_{(k)} - \widehat{\lambda}_{(k)}^{-1/2} \alpha b_{k2} \right\| + \left\| \widehat{\lambda}_{(k)}^{-1/2} \alpha (b_{(k)} - b_{k2}) \right\| \\ &\lesssim_P T^{-1/2} q^{-1/2} N^{-1/2} + q^{-1} N^{-1}. \end{aligned} \quad (37)$$

Also, with Lemma 3(iii), we have

$$|\widehat{\lambda}_{(k)}^{-1/2} - \lambda_{(k)}^{-1/2}| \leq \widehat{\lambda}_{(k)}^{-1/2} |\widehat{\lambda}_{(k)}^{1/2} / \lambda_{(k)}^{1/2} - 1| \lesssim_P T^{-1/2} q^{-1/2} N^{-1/2} + q^{-1} N^{-1}.$$

Since $\|b_{(k)}\| = 1$, the above two inequalities lead to

$$\left\| \widehat{\alpha}_{(k)} - \lambda_{(k)}^{-1/2} \alpha b_{(k)} \right\| \leq T^{-1/2} q^{-1/2} N^{-1/2} + q^{-1} N^{-1}. \quad (38)$$

For each term in the summation of (35), we have

$$\begin{aligned} &\left\| \widehat{\alpha}_{(k)} \zeta'_{(k)} \widetilde{D}_{(k)} u_T - \lambda_{(k)}^{-1/2} \alpha b_{(k)} \zeta'_{(k)} D_{(k)} u_T \right\| \\ &\leq \left\| \widehat{\alpha}_{(k)} (\zeta'_{(k)} \widetilde{D}_{(k)} u_T - \zeta'_{(k)} D_{(k)} u_T) \right\| + \left\| (\widehat{\alpha}_{(k)} - \lambda_{(k)}^{-1/2} \alpha b_{(k)}) \zeta'_{(k)} D_{(k)} u_T \right\|. \end{aligned} \quad (39)$$

Note that (37) also implies $\|\widehat{\alpha}_{(k)}\| \lesssim_P q^{-1/2} N^{-1/2}$ as $\widehat{\lambda}_{(k)} \asymp qN$, and that (36) implies the first term in (39) is $O_P(T^{-1} + q^{-1} N^{-1})$. Furthermore, $|\zeta'_{(k)} D_{(k)} u_T| \lesssim_P 1$ from Lemma 5(iv) and (38) show that the second term in (39) is also $O_P(T^{-1} + q^{-1} N^{-1})$. Given this, (35) becomes

$$\left\| \widehat{\gamma} u_T - \sum_{k=1}^K \lambda_{(k)}^{-1/2} \alpha b_{(k)} \zeta'_{(k)} D_{(k)} u_T \right\| \lesssim_P T^{-1} + q^{-1} N^{-1}. \quad (40)$$

To sum up, we have established that

$$\widehat{y}_{T+h} - E_T(y_{T+h}) = \frac{ZF'}{T_h} f_T + \frac{ZW'}{T_h} \Sigma_w^{-1} w_T + \sum_{k=1}^K \lambda_{(k)}^{-1/2} \alpha b_{(k)} \zeta'_{(k)} D_{(k)} u_T + O_P(T^{-1} + q^{-1} N^{-1}).$$

In the general case that Σ_f may not be $\mathbb{I}_K$, the first term becomes $T_h^{-1} ZF' \Sigma_f^{-1} f_T$. Using the fact $\varsigma_{(k)} = \lambda_{(k)}^{-1/2} \beta_{(k)} b_{(k)} = \lambda_{(k)}^{-1/2} \beta_{[I_k]} b_{(k)}$ and the iterative definition of $D_{(k)}$, we can see that $\lambda_{(k)}^{-1/2} \zeta'_{(k)} D_{(k)} u_T$ is exactly the k th row of $\Lambda^{-1} \Omega' \Psi u_T$ with Λ , Ω , and Ψ defined in Theorem 4. Using Delta method and Assumption 6, it is straightforward to obtain the desired CLT. $\square$

References

Ahn, S. C. and J. Bae (2022). Forecasting with partial least squares when a large number of predictors are available. Technical report, Arizona State University and University of Glasgow.

- Amini, A. A. and M. J. Wainwright (2009, October). High-dimensional analysis of semidefinite relaxations for sparse principal components. *Annals of Statistics* 37(5B), 2877–2921.
- Bai, J. (2003). Inferential Theory for Factor Models of Large Dimensions. *Econometrica* 71(1), 135–171.
- Bai, J. and S. Ng (2002). Determining the number of factors in approximate factor models. *Econometrica* 70, 191–221.
- Bai, J. and S. Ng (2006). Determining the number of factors in approximate factor models, erata. <http://www.columbia.edu/~sn2294/papers/correctionEcta2.pdf>.
- Bai, J. and S. Ng (2008). Forecasting economic time series using targeted predictors. *Journal of Econometrics* 146(2), 304–317.
- Bai, J. and S. Ng (2021). Approximate factor models with weaker loading. Technical report, Columbia University.
- Bai, Z. and J. W. Silverstein (2006). *Spectral Analysis of Large Dimensional Random Matrices*. Springer.
- Bai, Z. and J. W. Silverstein (2009). *Spectral Analysis of Large Dimensional Random Matrices*. Springer.
- Bair, E., T. Hastie, D. Paul, and R. Tibshirani (2006). Prediction by supervised principal components. *Journal of the American Statistical Association* 101(473), 119–137.
- Bair, E. and R. Tibshirani (2004). Semi-supervised methods to predict patient survival from gene expression data. *PLoS Biology* 2(4), 511–522.
- Cai, T. T., T. Jiang, and X. Li (2021). Asymptotic analysis for extreme eigenvalues of principal minors of random matrices. *The Annals of Applied Probability* 31(6), 2953–2990.
- Chamberlain, G. and M. Rothschild (1983). Arbitrage, factor structure, and mean-variance analysis on large asset markets. *Econometrica* 51, 1281–1304.
- Chao, J. C. and N. R. Swanson (2022). Consistent estimation, variable selection, and forecasting in factor-augmented var models. Technical report, University of Maryland and Rutgers University.
- d’Aspremont, A., L. E. Ghaoui, M. I. Jordan, and G. R. G. Lanckriet (2007, January). A direct formulation for sparse PCA using semidefinite programming. *SIAM Review* 49(3), 434–448.
- Fan, J., Y. Liao, and M. Mincheva (2011). High-dimensional covariance matrix estimation in approximate factor models. *Annals of Statistics* 39(6), 3320–3356.
- Forni, M., D. Giannone, M. Lippi, and L. Reichlin (2009). Opening the black box: Structural factor models with large cross sections. *Econometric Theory* 25, 1319–1347.
- Forni, M., M. Hallin, M. Lippi, and L. Reichlin (2000). The generalized dynamic-factor model: Identification and estimation. *The Review of Economics and Statistics* 82, 540–554.
- Forni, M., M. Hallin, M. Lippi, and L. Reichlin (2004, April). The generalized dynamic factor model:

- Consistency and rates. *Journal of Econometrics* 119(2), 231–255.
- Forni, M. and M. Lippi (2001). The generalized dynamic factor model: Representation theory. *Econometric Theory* 17, 1113–1141.
- Giglio, S., D. Xiu, and D. Zhang (2020). Test assets and weak factors. Technical report, Yale University and University of Chicago.
- Hoyle, D. C. and M. Rattray (2004). Principal-component-analysis eigenvalue spectra from data with symmetry-breaking structure. *Physical Review E* 69(2), 026124.
- Huang, D., F. Jiang, K. Li, G. Tong, and G. Zhou (2021). Scaled pca: A new approach to dimension reduction. *Management Science*, *forthcoming*.
- Johnstone, I. M. (2001). On the distribution of the largest eigenvalue in principal components analysis. *Annals of Statistics* 29, 295–327.
- Johnstone, I. M. and A. Y. Lu (2009). On consistency and sparsity for principal components analysis in high dimensions. *Journal of the American Statistical Association* 104(486), 682–693.
- Jolliffe, I. T., N. T. Trendafilov, and M. Uddin (2003, September). A modified principal component technique based on the LASSO. *Journal of Computational and Graphical Statistics* 12(3), 531–547.
- Kelly, B. and S. Pruitt (2013). Market expectations in the cross-section of present values. *The Journal of Finance* 68(5), 1721–1756.
- Onatski, A. (2009). Testing hypotheses about the number of factors in large factor models. *Econometrica* 77(5), 1447–1479.
- Onatski, A. (2010). Determining the number of factors from empirical distribution of eigenvalues. *Review of Economics and Statistics* 92, 1004–1016.
- Onatski, A. (2012). Asymptotics of the principal components estimator of large factor models with weakly influential factors. *Journal of Econometrics* 168, 244–258.
- Paul, D. (2007). Asymptotics of sample eigenstructure for a large dimensional spiked covariance model. *Statistical Sinica* 17, 1617–1642.
- Stock, J. H. and M. W. Watson (2002). Forecasting Using Principal Components from a Large Number of Predictors. *Journal of the American Statistical Association* 97(460), 1167–1179.
- Uematsu, Y., Y. Fan, K. Chen, J. Lv, and W. Lin (2019). Sofar: Large-scale association network learning. *IEEE Transactions on Information Theory* 65(8), 4924–4939.
- Uematsu, Y. and T. Yamagata (2021). Estimation of sparsity-induced weak factor models. *Journal of Business and Economic Statistics*, *forthcoming*.
- Wang, W. and J. Fan (2017). Asymptotics of empirical eigenstructure for ultra-high dimensional spiked covariance model. *Annals of Statistics* 45, 1342–1374.
- Zou, H., T. Hastie, and R. Tibshirani (2006). Sparse principal component analysis. *Journal of Computational and Graphical Statistics* 15, 265–286.

Online Appendix of

Prediction When Factors are Weak

Stefano Giglio*

Yale School of Management

NBER and CEPR

Dacheng Xiu[†]

Booth School of Business

University of Chicago

Dake Zhang[‡]

Booth School of Business

University of Chicago

This Version: March 28, 2023

Abstract

The online appendix presents an empirical application, details on the asymptotic variance estimation, and proofs of propositions and technical lemmas.

*Address: 165 Whitney Avenue, New Haven, CT 06520, USA. E-mail address: stefano.giglio@yale.edu.

[†]Address: 5807 S Woodlawn Avenue, Chicago, IL 60637, USA. E-mail address: dacheng.xiu@chicagobooth.edu.

[‡]Address: 5807 S Woodlawn Avenue, Chicago, IL 60637, USA. Email: dkzhang@chicagobooth.edu.

A Empirical Analysis

In this section we apply our methodology in a standard macroeconomic prediction exercise, using a large set of predictors to forecast inflation, industrial production, and unemployment.

A.1 Empirical Context

Predicting macroeconomic variables like output and inflation is a central exercise in empirical macroeconomics. The availability of large macroeconomic datasets that contain many potentially useful predictors has spurred the application of a variety of methods of dimension reduction to this objective. Some of these methods, like those based on principal component analysis (PCA), reduce the dimensionality of the predictors universe without using information in the target of the forecast (see [Stock and Watson \(2002\)](#)). Others instead use information from the target to help the dimension reduction focus on the most valuable predictors; examples include partial least squares (PLS, [Kelly and Pruitt \(2015\)](#)), targeted PCA ([Bai and Ng \(2008\)](#)), and scaled PCA ([Huang et al. \(2022\)](#)). SPCA belongs to the latter group, as it employs an iterative screening step based on correlation with the target to eliminate useless or noisy predictors.

Because the selection step is designed to eliminate irrelevant predictors (as opposed to downweight them as, for example, PLS does) we expect SPCA to perform best when faced with a large number of predictors that are potentially irrelevant, noisy, or redundant. In our empirical analysis, we therefore explore a context in which a large number of predictors are available to be used for forecasting. Specifically, we include in our set of predictors not only a standard panel of macroeconomic variables, but also a large dataset of individual forecasts of different macroeconomic quantities by professional forecasters. Macroeconomic forecasts have often been included in forecasting exercises, either by using the consensus forecast as an additional predictor ([Faust and Wright \(2013\)](#)) or in the context of optimal forecast combination ([Genre et al. \(2013\)](#)). In our context, we let SPCA decide if and which individual forecasts to use to complement the macroeconomic predictors – so the forecast combination will be decided automatically by SPCA.

A.2 Data

Our empirical exercise combines two datasets. First, we use the standard Fred-Md database ([McCracken and Ng, 2016](#)) that contains 127 monthly macroeconomic and financial series.¹ The Fred-Md

¹The series are grouped in the following categories: output and income; labor market; housing; consumption, orders and inventories; money and credit; interest and exchange rates; prices; stock market. The dataset applies a variety of transformations to the underlying series, which we follow in our analysis. We however make a few adjustments to the series' data transformations, to ensure that all series are stationary and based on economic reasoning. For the Effective Federal Funds Rate (FEDFUNDS), we keep its level (i.e., no transformation) instead of taking the first difference. We also compute the first difference of natural log instead of the second difference of natural log for the following series: M1 Money Stock (M1SL), M2 Money Stock (M2SL), Board of Governors Monetary Base (BOGMBASE; note: starting from the January 2020 (2020-01) vintage, BOGMBASE replaced the St. Louis Adjusted Monetary Base (AMBSL)), Total Reserves of Depository Institutions (TOTRESNS), Commercial and Industrial Loans (BUSLOANS), Real Estate Loans

data spans the period March 1959 to February 2022. Second, we use individual forecasts from the Blue Chip Financial Forecasts data, which is a monthly survey of experts from various major financial institutions³ and provides forecasts of interest rates and many other macroeconomic quantities⁴ for each of the next six quarters (i.e., current quarter t through $t + 5$), for a total of hundreds of forecasts every month. Our data covers the period February 1993 to February 2022 and we use all forecasts available (for all possible macroeconomic targets) as potential predictors. This gives us up to 18,053 different individual forecasts that could in theory be used as predictors (though, as discussed below, many of these forecasts are available for only a small number of periods, so they are not used in our analysis). Given that the Blue Chip forecast is only available since 1993, we conduct all of our analysis for the period February 1993 to February 2022.

A.3 Out of Sample Forecast Evaluation

We forecast each of the three targets (inflation, industrial production growth, and change in the unemployment rate) using a rolling out of sample procedure. We evaluate the out of sample forecast of SPCA and compare it with two alternative forecasting methods, PCA and PLS. We choose these alternatives because each is a prominent example of a class of methods used in large-dimensional macroeconomic forecasting (respectively, unsupervised and supervised dimension reduction). Each of the three methods we evaluate (SPCA, PCA, PLS) is benchmarked to the forecast of an autoregressive model, whose number of lags is selected by the BIC criterion with a maximum lag of 12 lags, using a direct projection approach (Marcellino et al. (2006), Faust and Wright (2013)). We study forecast horizons of 1 to 12 months.

All of the analysis is performed using a rolling estimation on a 240-months window. At every time t starting at the last month of the window, we predict the cumulated macroeconomic variables from t to $t + h$, where h is the forecast horizon, as in Huang et al. (2022). Within each window, we only keep predictors that have less than 10% missing data points. For those series that are included but

at All Commercial Banks (REALLN), Total Nonrevolving Credit (NONREVSL), Finished Goods (WPSFD49207), Finished Consumer Goods (WPSFD49502), Processed Goods for Intermediate Demand (WPSID61), Unprocessed Goods for Intermediate Demand (WPSID62; note: starting from the March 2016 (2016-03) vintage, PPI: Finished Goods (PPIFGS), PPI: Finished Consumer Goods (PPIFCG), PPI: Intermediate Materials (PPIITM), and PPI: Crude Materials (PPICRM) have been replaced with WPSFD49207, WPSFD49502, WPSID61, and WPSID62 respectively), Crude Oil, spliced WTI and Cushing (OILPRICE), PPI: Metals and Metal Products (PPICMM), Consumer Price Index for All Urban Consumers (CPIAUCSL), CPI: Apparel (CPIAPPSL), CPI: Transportation (CPITRNSL), CPI: Medical Care (CPIMEDSL), CPI: Commodities (CUSR0000SAC), CPI: Durables (CUSR0000SAD), CPI: Services (CUSR0000SAS), CPI: All Items Less Food (CPIULFSL), CPI: All Items Less Shelter (CUSR0000SA0L2)², CPI: All Items Less Medical Care (CUSR0000SA0L5), Personal Cons. Exp: Chain Index (PCEPI), Personal Cons. Exp: Durable Goods (DDURRG3M086SBEA), Personal Cons. Exp: Nondurable Goods (DNDGRG3M086SBEA), Personal Cons. Exp: Services (DSERRG3M086SBEA), Avg Hourly Earnings: Goods-Producing (CES0600000008), Avg Hourly Earnings: Construction (CES2000000008), Avg Hourly Earnings: Manufacturing (CES3000000008), Consumer Motor Vehicle Loans Outstanding (DTCOLNVHFN), Total Consumer Loans and Leases Outstanding (DTCTHFN) and Securities in Bank Credit at All Commercial Banks (INVEST).

³For instance, Bank of America, Goldman Sachs & Co. and J.P. MorganChase.

⁴For instance, the percentage changes in Real GDP, the GDP Chained Price Index, the Consumer Price Index and a set of interest rates (e.g., Federal Funds, 3-month Treasury, Aaa as well as Baa Corporate Bonds).

do have some missing data (mostly Blue Chip forecasts) we forward fill the last non-missing value. About half of the total of around 40 forecasters from BlueChip available in the average month have sufficiently long series of forecasts to be included in our analysis. All predictors are standardized within each window. Then, a forecast is made for $t + 1$ using the three different methods, and these forecasts are then joined over time to compute the out-of-sample R^2 (relative to the AR benchmark). When we use the Blue Chip data, we also include dummies for month of the quarter, to account for the fact that the Blue Chip data makes forecasts for calendar quarters irrespective of the month.⁵

Recall that the SPCA procedure presented in Section 2 relies on two tuning parameters, K and $\lfloor qN \rfloor$, whereas PCA and PLS only rely on tuning K . To demonstrate the effect of tuning parameters, we report three versions of the results. We first show the performance of the forecasting methods for different (fixed) number of factors K and different (fixed) choice of $\lfloor qN \rfloor$. In this case, no tuning is needed for SPCA. We then show the performance of SPCA for each K , with a single tuning parameter of SPCA that drives the selection step $\lfloor qN \rfloor$ chosen via 3-fold cross-validation (CV) separately in each time window. Next, we show the results when both the number of factors K (for SPCA, PCA and PLS) and the tuning parameter $\lfloor qN \rfloor$ (for SPCA) are jointly chosen via CV. We consider a range of $\lfloor qN \rfloor$ from 50 to 300.

A.4 Results

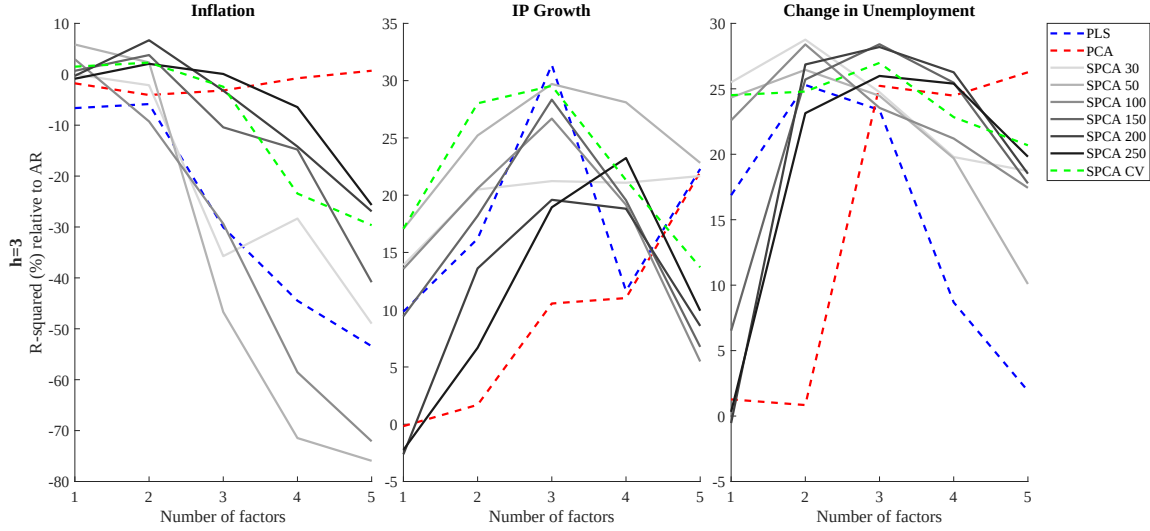
A.4.1 Forecasting Performance

We begin by focusing on prediction at the quarterly (3-month) horizon, which is a standard horizon studied in the literature. Figure A1 reports the out of sample R^2 of different forecasting methodologies relative to the AR benchmark, for inflation (left panel), industrial production growth (center panel), and change in unemployment (right panel). In this figure, the prediction exercise is performed by fixing the number of factors K . For PCA (red line) and PLS (blue line), there are no tuning parameters beyond K . For SPCA, we report separate results for each choice of the tuning parameter K (grey lines), as well as for the value for $\lfloor qN \rfloor$ chosen by CV (green line).

The figure shows several interesting results. First, it is in general hard to predict inflation beyond what an AR model predicts (see also Faust and Wright (2013)): the out of sample R^2 s are close to zero or even negative. Only SPCA, among all methods, produces positive R^2 s, and it does so using a small number of factors. Predictability beyond the AR model is much higher for IP growth and change in unemployment. Second, the predictive performance of SPCA is generally higher than that of PCA and PLS for most choices of the number of factors. Third, the performance of PCA does depend on the tuning parameter, but in different ways for different targets. For inflation, for example, a lower value of $\lfloor qN \rfloor$ seems to predict better; for industrial production and unemployment, higher values work better. Finally, the performance of all these methods varies quite dramatically with the number of factors, with substantial declines for the methods that use target information

⁵For example, in January, February and March, the “current quarter” forecast always refers to Q1.

Figure A1: OOS Performance of SPCA, PCA and PLS (for different number of factors)



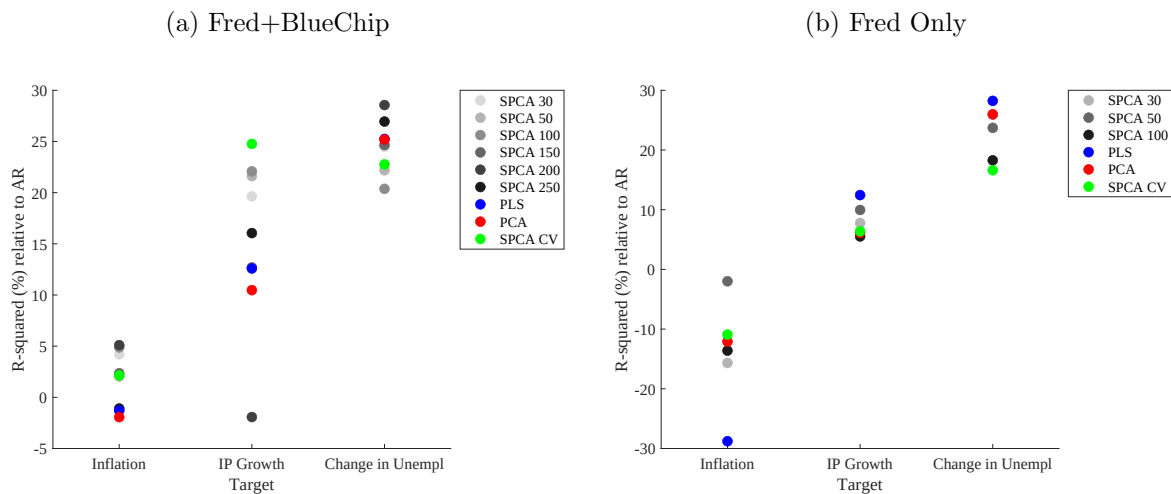
Notes: Each panel reports the out-of-sample R^2 relative to the AR model for a different target, aggregated over 3 months. The three panels predict inflation, industrial production growth and change in unemployment rate, respectively. The green dashed line shows the performance of SPCA with 3-fold cross validation for the tuning parameter $[qN]$. The grey lines show the performance of SPCA with fixed number of predictors, $[qN]$. The blue dashed line uses PLS. The red dashed line uses PCA. Rolling window of 240 months is used. Sample covers 1993-2022.

(PLS and SPCA with a smaller $[qN]$) as the number of factors increases, because of their overfitting issue we explained earlier.

Given how important the number of factors is for the out-of-sample performance, in what follows we choose the number of factors via cross-validation for all three methods (so for SPCA both $[qN]$ and K are jointly selected via CV). The left panel of Figure A2 shows the results. Now all three targets (inflation, industrial production growth and change in unemployment rate) appear in the same panel. The panel confirms that SPCA generally performs well in predicting out of sample, doing better than the alternatives (in the case of unemployment, several choices of the tuning parameter $[qN]$ outperform PCA and PLS, but not the one chosen by cross-validation). Overall, SPCA tends to do comparatively well when choosing all parameters via cross-validation.

Given the way SPCA chooses the set of predictors, we would expect it to perform best in contexts where there are a large number of predictors, that overall contain valuable information, even if some predictors are redundant or noisy. The forecasting experiment we run here falls in this category: it contains both macroeconomic and financial data (which are likely to contain important individual predictors), as well as a large number of individual forecasts that we would expect to be informative beyond macroeconomic quantities but where a large part of the observed variation is likely dominated by noise. To better gauge the importance of this additional data in the performance of SPCA, the right panel of Figure A2 shows the results of running the same analysis (using the same sample) but

Figure A2: OOS Performance of SPCA, PCA and PLS (using CV to choose the number of factors)



Notes: The left panel of this figure repeats the analysis of Figure A1, but chooses the number of factors via CV. The right panel performs the same analysis as the left panel, but using only Fred data.

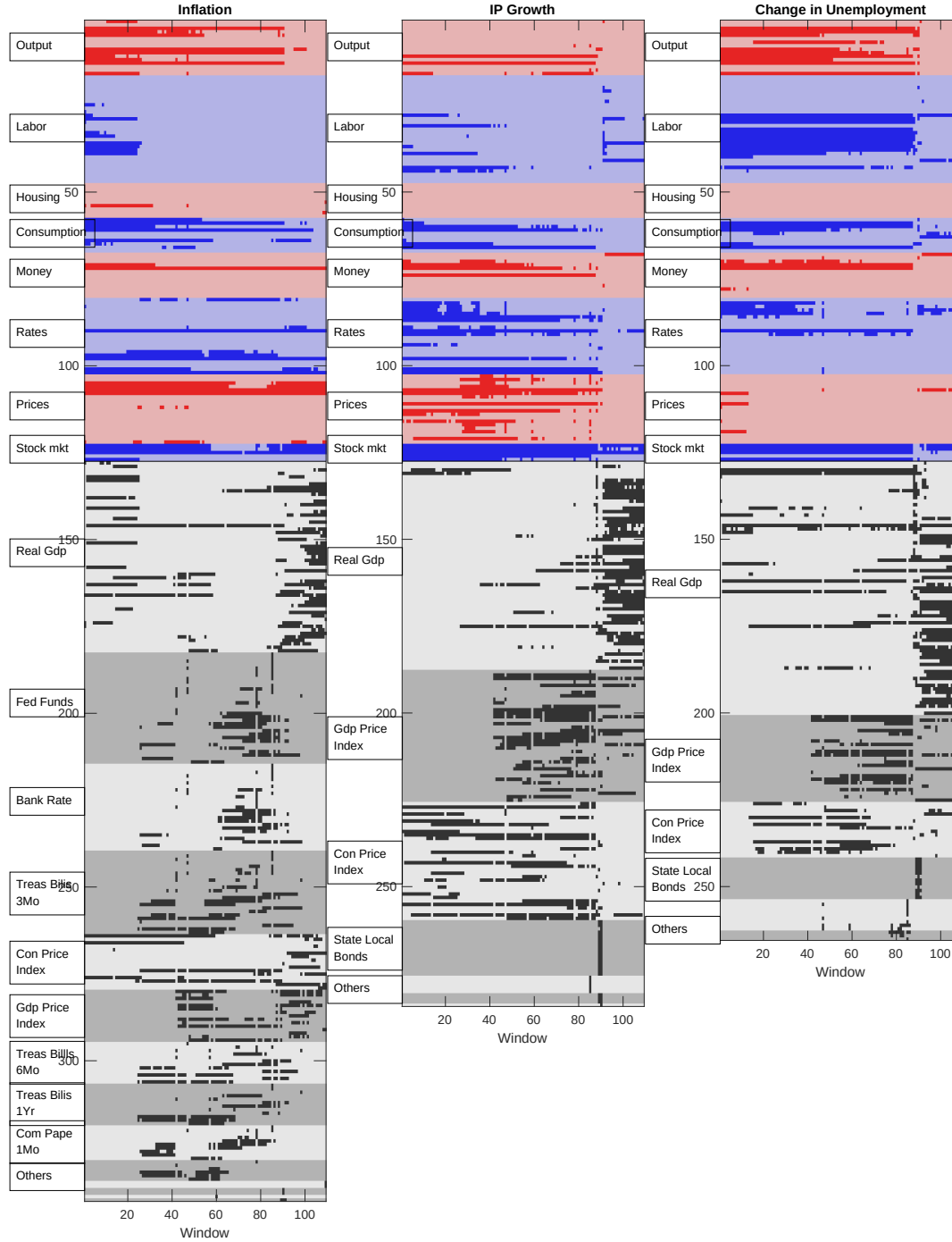
with only the Fred data. The figure shows that while the performance of SPCA remains broadly comparable with the other predictors, it deteriorates compared to PCA and PLS (PLS itself has very mixed performance, though, predicting well IP growth and unemployment, and failing to predict inflation). So, on the one hand, this figure shows that individual expert forecasts are useful for prediction of macroeconomic variables, confirming the results in Faust and Wright (2013); on the other hand, it shows that SPCA does particularly well when working with this large and informative, yet noisy, universe of individual forecasts.

A.4.2 Predictors Selected by SPCA

Next, we study in detail how SPCA selects predictors. Figure A3 shows which variables are chosen by SPCA to extract the first factor (focusing on the 50 with highest correlation with the target, for reasons of readability). For the three targets (one per column), the graph reports which variables were selected in each of the rolling windows in our sample. The top part of the graph collects the 127 Fred variables, grouped according to the standard Fred-Md categorization, in alternating blue and red colors. The bottom part corresponds to the BlueChip surveys, grouped by the target of the individual forecast (therefore, each row in this part of the graph is a forecast of a particular variable, at a particular horizon, by a particular expert). A darker color in this graph means that the variable is selected in that window.

Consider for example the inflation graph on the left. To extract a factor useful to predict inflation, SPCA selects a large number of variables from a few groups: output, consumption, rates, prices, and the stock market. Other groups are almost never selected. Rates are selected more for IP growth, and labor variables are selected more when predicting unemployment. Housing variables are rarely used for all three targets. Note that in many cases, the same predictors from each group are used,

Figure A3: Top 50 Predictors Selected by SPCA



Notes: Under the same settings as Figure A1, each panel visualizes the top 50 predictors selected by SPCA across windows while predicting each target. The first set of variables (in red and blue) are Fred predictors, and the second set (in grey) corresponds to the BlueChip forecasts. For the latter set, only the predictors ever among the top 50 by correlation with the target are visualized.

indicating that the predictive power of these macroeconomic variables is persistent.

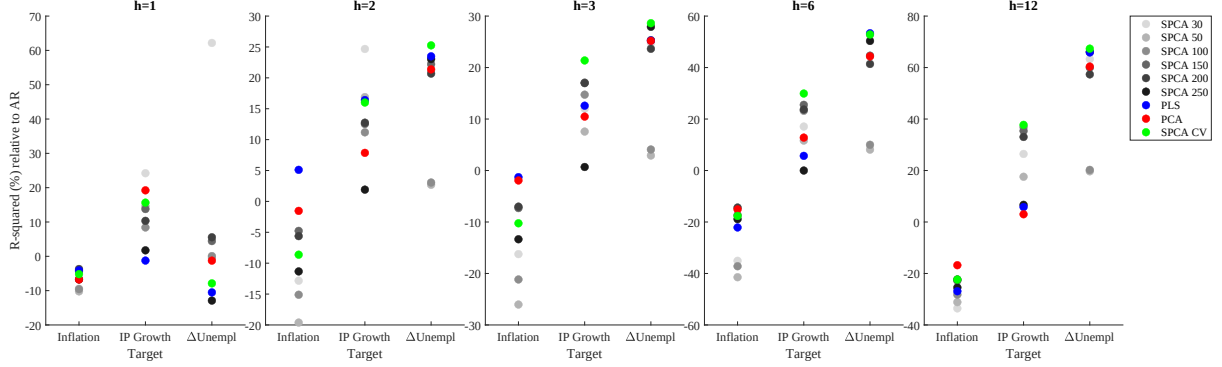
To this macroeconomic set of predictors, SPCA adds a selection of individual forecasts from the BlueChip data as additional predictors. For reasons of space, the greyscale part of the graph shows a subset of these predictors: only those that are selected among the top 50 predictors at least in one window. The graph shows that different types of forecasts are used at different points in time, with some exceptions. Not surprisingly, to predict inflation, forecasts of the consumer price index are always included. To these forecasts, SPCA adds forecasts of GDP in the first and last part of the sample, and interest rates in the intermediate part of the sample. GDP forecasts are used throughout the sample to predict changes in unemployment, and become more dominant for all target variables toward the end of the sample, whereas inflation predictors tend to be more important beforehand. This switch is perhaps due to the fact that in the later part of the sample the zero lower bound was close or binding and inflation was low and not very volatile.

Finally, we note that not all Blue Chip forecasters are the same in terms of forecasting ability. Among the institutions whose forecasts are included in our analysis because they have a sufficiently long time series (each providing tens of forecasts, of different variables at different horizons), we find significant heterogeneity in the frequency with which their forecasts are selected by SPCA. For example, Nomura has its forecasts selected between 23% and 39% of the time at the first iteration (depending on the target). Swiss RE, on the other hand, has its forecasts selected only 0.1% of the time, for each target. This distribution is quite skewed: only 5 institutions have their forecasts selected more than 10% of the time for each target, out of the 20 included in our sample. Similar results hold when looking at selection at any iteration of SPCA.

A.4.3 Joint Forecasts using Many Targets

Next, one special feature of SPCA is that it can operate the selection using a set of multiple targets jointly. In fact, using multiple targets is required by the theory (see Section 3.2) to do inference, as long as there are more than one factors in the true DGP. We implement this here by predicting each target at horizons of 1, 2, 3, 6 and 12 months jointly. Figure A4 reports the out of sample R^2 s on each horizon. There are two main results that this figure highlights. First, SPCA tends to do on average well at longer horizons (3, 6 and 12 months), whereas its performance is more uneven at shorter horizons. Second, comparing the middle panel (predicting one quarter ahead) with the left panel of Figure A2, which focused on the 3-month horizon only, we see that the use of other horizons to help select predictors has different effects for different targets. It significantly improves the forecasting ability for unemployment, but reduces the forecasting ability for IP growth (mildly) and inflation (significantly so). Overall, the performance of SPCA remains on par with the other predictors when using multiple targets, especially at longer horizons.

Figure A4: OOS Performances - Different Targeted Horizons



Notes: Similar to Figure A2, but showing the out of sample R^2 s at different horizons, and using all the horizons concurrently to estimate the factors in SPCA.

A.4.4 Time Series of the Forecasts

Finally, we study the time series of our out-of-sample forecasts at different horizons, using the estimates obtained in Section A.4.3, for horizons of 1, 2, 3, 6 and 12 months. Figure A5 reports SPCA's forecasts with asymptotic forecast standard errors at each maturity. In the figure, the blue dots represent the underlying time series that is the target of the forecast: log CPI, log IP, and unemployment, all scaled to start from 0 at the beginning of the sample. For readability, we show the forecasts every six months, each for horizons up to 12 months. Standard errors are obtained using the asymptotic distributions derived in Section 3.2, and are plotted in three shades (the 10th and 90th percentiles in the darkest shade, 5th and 95th in the middle shade, and 1st and 99th in the lightest shade).

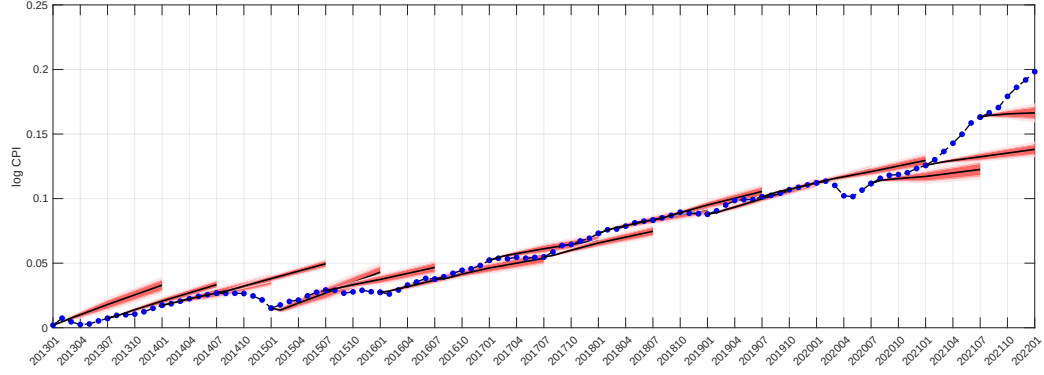
Overall, SPCA does a good job forecasting the three series, with the forecasts often anticipating changes in the direction of the different variables. For example, IP forecasts predicted the increase starting in 2016, and the decrease that started in 2018. Of course, in other times the forecasts miss significantly, sometimes for several periods in the same direction. Two examples: first, forecasts do not fully anticipate the persistent decrease in unemployment that occurred during 2013 and 2014. Second, all forecasts miss (as they should have) the unexpected and extraordinary events of the Covid pandemic (both the initial shock and the recovery). In that period, the point estimates change dramatically over a short period of time, and standard errors increase noticeably, demonstrating the large amount of uncertainty about the path of the economy during those times.

B Estimation of Φ_1 and Φ_2

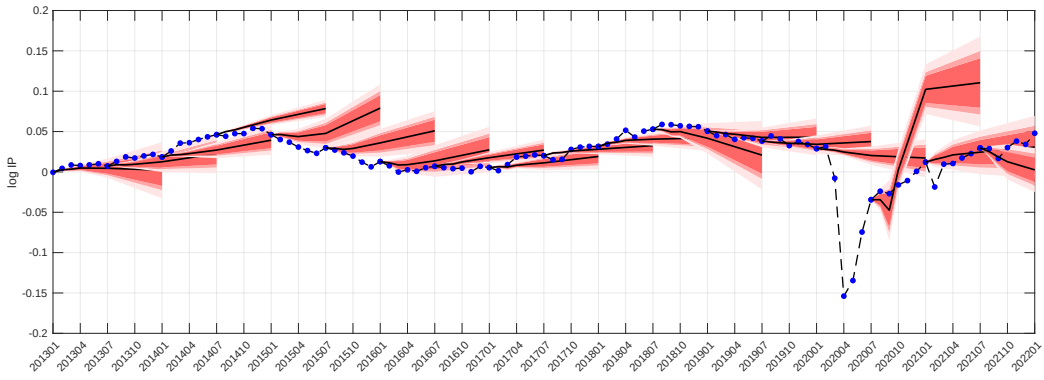
Recall that from the outputs of Algorithm 1, we have defined $\hat{F}$, $\hat{\beta}$, and $\hat{\alpha}$. As a result, we can also estimate $\hat{Z} = Y - \hat{\alpha}\hat{F} - \hat{\alpha}_w W$ and $\hat{U} = X - \hat{\beta}\hat{F} - \hat{\beta}_w W$. Then we can construct Newey-West-type estimators for Π_{11} , Π_{12} and Π_{22} , given that each component of them can be estimated based on their

Figure A5: Fan Charts

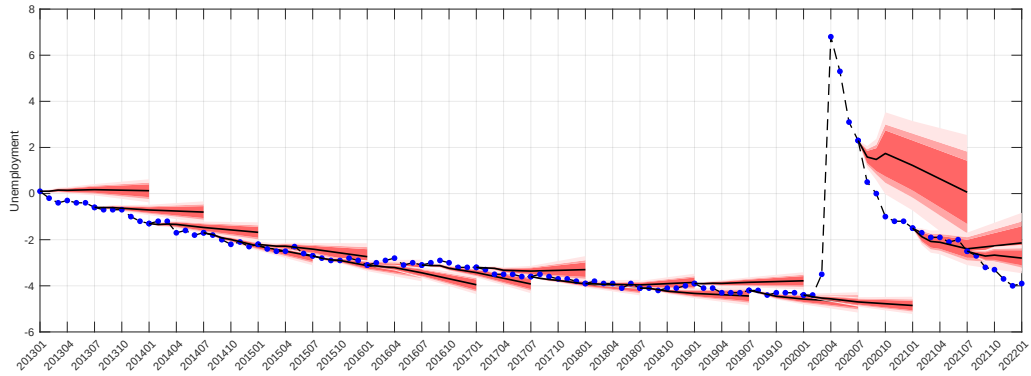
(a) Inflation



(b) IP Growth



(c) Change in Unemployment



Notes: Using the same estimates as Figure A4, each panel shows the forecasts and confidence intervals for horizons up to 12 months. The forecasts are shown every 6 months, in alternating red and green colors (for readability). The blue dots are the cumulative targets of the forecasts.

sample analog constructed above. Estimators of Σ_f and Σ_w can be obtained by $\widehat{\Sigma}_f = T_h^{-1} \widehat{F} \widehat{F}'$ and $\widehat{\Sigma}_w = T_h^{-1} \widehat{W} \widehat{W}'$. With $\widehat{f}_T = \widehat{\zeta}'(x_T - \widehat{\beta}_w w_T)$, $\widehat{\Phi}_1$ can be constructed as follows:

$$\widehat{\Phi}_1 = \left((\widehat{f}_T', w_T') \widehat{\Sigma}_{f,w}^{-1} \otimes \mathbb{I}_D \right) \begin{pmatrix} \widehat{\Pi}_{11} & \widehat{\Pi}_{12} \\ \widehat{\Pi}_{12}' & \widehat{\Pi}_{22} \end{pmatrix} \left(\widehat{\Sigma}_{f,w}^{-1} (\widehat{f}_T', w_T')' \otimes \mathbb{I}_D \right).$$

The above estimators are built as if the latent factors were observed. This is because any rotation matrix involved with latent factor estimates is canceled out, which eventually yields consistent estimators of Φ_1 . This part of the asymptotic variance is straightforward to implement, thanks to the fact that it does not involve estimation of high-dimensional quantities like Σ_u . The proof of consistency of $\widehat{\Phi}_1$ follows directly from [Giglio and Xiu \(2021\)](#) and is thus omitted here.

With respect to Φ_2 , we may apply a thresholding estimator of $\Sigma_u = \text{Cov}(u_t)$ following [Fan et al. \(2013\)](#). In detail, $\widehat{\Sigma}_u$ can be constructed by

$$(\widehat{\Sigma}_u)_{ij} = \begin{cases} (\widetilde{\Sigma}_u)_{ij}, & i = j \\ s_{ij} \left((\widetilde{\Sigma}_u)_{ij} \right), & i \neq j \end{cases}, \quad \widetilde{\Sigma}_u = T_h^{-1} \widehat{U} \widehat{U}',$$

where $s_{ij}(\cdot)$ is a general thresholding function with an entry-dependent threshold τ_{ij} satisfying (i) $s_{ij}(z) = 0$ when $|z| \leq \tau_{ij}$ (ii) $|s_{ij}(z) - z| \leq \tau_{ij}$. The adaptive threshold can be chosen by $\tau_{ij} = C \left(\frac{1}{\sqrt{qN}} + \sqrt{\frac{\log N}{T}} \right) \sqrt{\widehat{\theta}_{ij}}$, where $C > 0$ is a sufficiently large constant and

$$\widehat{\theta}_{ij} = \frac{1}{T_h} \sum_{t \leq T_h} (\widehat{u}_{it} \widehat{u}_{jt} - (\widetilde{\Sigma}_u)_{ij})^2,$$

where $\widehat{u}_{it}$ are the entries of $\widehat{U}$. With $\widehat{\Sigma}_u$, Φ_2 can be estimated by $\widehat{\Phi}_2 = qN \widehat{\gamma} \widehat{\Sigma}_u \widehat{\gamma}'$.

The following theorem ensures the consistency of $\widehat{\Phi}_2$ under standard assumptions as in [Fan et al. \(2013\)](#).

Theorem B.1. *Under the assumptions of Theorem 4, if we further assume that*

- (i) u_t is stationary with $E(u_t) = 0$ and $\Sigma_u = \text{Cov}(u_t)$ satisfying $C_1 > \lambda_1(\Sigma_u) \geq \lambda_N(\Sigma_u) > C_2$ and $\min_{i,j} \text{Var}(u_{it} u_{jt}) > C_2$ for some constant $C_1, C_2 > 0$,
- (ii) u_t has exponential tail, i.e., there exist $r_1 > 0$ and $C > 0$, such that for any $s > 0$ and $i \leq N$, $P(|u_{it}| > s) \leq \exp(-(s/C)^{r_1})$.
- (iii) u_t is strong mixing, i.e., there exist positive constants r_2 and C such that for all $t \in \mathbb{Z}^+$, $\alpha(t) \leq \exp(-Ct^{r_2})$, where

$$\alpha(T) = \sup_{A \in \mathcal{F}_{-\infty}^0, B \in \mathcal{F}_T^\infty} |P(A)P(B) - P(AB)|$$

and $\mathcal{F}_{-\infty}^0, \mathcal{F}_T^\infty$ are σ -algebras generated by $\{u_t\}_{-\infty \leq t \leq 0}, \{u_t\}_{T \leq t \leq \infty}$.

(iii) $(\log N)^{6(3r_1^{-1}+r_2^{-1}+1)} = o(T), T = o(q^2 N^2)$.

Then $\widehat{\Sigma}_u$ defined above satisfies $\|\widehat{\Sigma}_u - \Sigma_u\| \lesssim_P m_{q,N} \left(\frac{1}{\sqrt{qN}} + \sqrt{\frac{\log N}{T}} \right)^{1-q}$, where $m_{q,N} = \max_{i \leq N} \sum_{j \leq N} |(\Sigma_u)_{ij}|^q$. In addition, if $m_{q,N} \left(\frac{1}{\sqrt{qN}} + \sqrt{\frac{\log N}{T}} \right)^{1-q} = o(1)$, then $\widehat{\Phi}_2 \xrightarrow{P} \Phi_2$.

C Additional Mathematical Proofs

C.1 Proof of Theorem B.1

Proof. Using Theorem 5 in Fan et al. (2013), to establish the error bound $\|\widehat{\Sigma}_u - \Sigma_u\|$, it is sufficient to show that

$$\|\widehat{U} - U\|_{\text{MAX}} = o_P(1)$$

and

$$\max_{i \leq N} T_h^{-1} \sum_t |u_{it} - \widehat{u}_{it}|^2 = O_P \left(\frac{1}{qN} + \frac{\log N}{T} \right).$$

These two estimates have been shown by Lemma 11(iii)(iv). If $m_{q,N} \left(\frac{1}{\sqrt{qN}} + \sqrt{\frac{\log N}{T}} \right)^{1-q} = o(1)$, then $\|\widehat{\Sigma}_u - \Sigma_u\| = o_P(1)$. With $\|\zeta'_{(k)} \widetilde{D}_{(k)} - \varsigma'_{(k)} D_{(k)}\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$ from Lemma 8(iv) and $\widehat{\gamma} = \sum_{k \leq K} \widehat{\alpha}_{(k)} \zeta'_{(k)} \widetilde{D}_{(k)}$, rewrite the proof of (40), we have

$$\left\| \widehat{\gamma} - \sum_{k \leq K} \lambda_{(k)}^{-1/2} \alpha b_{(k)} \varsigma'_{(k)} D_{(k)} \right\| \lesssim_P T^{-1/2} q^{-1/2} N^{-1/2} + q^{-1} N^{-1}. \quad (\text{C.1})$$

The difference between the rate of (40) and this equation arises from the difference between Lemma 8(iv) and (vi). Recall that $\lambda_{(k)}^{-1/2} \varsigma'_{(k)} D_{(k)}$ is exactly the k th row of $\Lambda^{-1} \Omega' \Psi$, the left hand side of (C.1) is equivalent to $\|\widehat{\gamma} - \alpha B \Lambda^{-1} \Omega' \Psi\|$. In addition, under the assumption $\text{Cov}(u_t) = \Sigma_u$, Π_{33} equals to $(qN)^{-1} \Psi \Sigma_u \Psi'$. Let $\tilde{\gamma}$ denote $\alpha B \Lambda^{-1} \Omega' \Psi$, then we have

$$\widehat{\Phi}_2 - \Phi_2 = qN \left(\widehat{\gamma} \widehat{\Sigma}_u \widehat{\gamma}' - \tilde{\gamma} \Sigma_u \tilde{\gamma}' \right).$$

Consequently, we have

$$\|\widehat{\Phi}_2 - \Phi_2\| \leq qN \left(\|\widehat{\gamma}(\widehat{\Sigma}_u - \Sigma_u)\widehat{\gamma}'\| + \|(\widehat{\gamma} - \tilde{\gamma})\Sigma_u \widehat{\gamma}'\| + \|\tilde{\gamma}\Sigma_u(\widehat{\gamma} - \tilde{\gamma})'\| \right). \quad (\text{C.2})$$

Using the definition of $D_{(k)}$, $\|\beta_{[I_k]}\| \lesssim (qN)^{1/2}$, and $\lambda_{(k)} \asymp qN$, we have $\|D_{(k)}\| \lesssim 1$ and thus $\|\hat{\gamma}\| \lesssim q^{-1/2}N^{-1/2}$. Using $\|\hat{\Sigma}_u - \Sigma_u\| = o_P(1)$, (C.1), $\|\Sigma_u\| \lesssim 1$ from the assumption and $\|\hat{\gamma}\| \lesssim q^{-1/2}N^{-1/2}$, it is straightforward to verify that all three terms in (C.2) are $o_P(1)$. This completes the proof. $\square$

C.2 Proof of Propositions 1 and 2

Proof. Note that for any orthogonal matrix $\Gamma \in \mathbb{R}^{N \times N}$, the estimators based on PCA and PLS on ΓR are the same as those based on R . Thus, without loss of generality, we can assume $\beta = (\lambda^{1/2}, 0, \dots, 0)'$, where $\lambda = \|\beta\|^2$ and it will not affect A .

We can then write X in the following form:

$$X = \beta F + U = \beta F + \epsilon A_1 = \begin{pmatrix} \sqrt{\lambda}F + \epsilon_1 A_1 \\ \epsilon_2 A_1 \end{pmatrix}, \quad (\text{C.3})$$

where ϵ_1 is the first row of ϵ and ϵ_2 contains the remaining rows. Correspondingly, we write the first left singular vector of X as $\hat{\varsigma} = (\hat{\varsigma}_1, \hat{\varsigma}_2)'$, where $\hat{\varsigma}_1$ is the first element of $\hat{\varsigma}$ and $\hat{\varsigma}_2$ is a vector of the remaining $N - 1$ entries of $\hat{\varsigma}$, write $\hat{\xi}$ as the first right singular vector of X , and denote the first singular value as $\sqrt{T\lambda}$. By simple algebra we have

$$\hat{\varsigma}_1 = \frac{(\sqrt{\lambda}F + \epsilon_1 A_1)\hat{\xi}}{\sqrt{T\lambda}}, \quad \hat{\varsigma}_2 = \frac{\epsilon_2 A_1 \hat{\xi}}{\sqrt{T\lambda}}. \quad (\text{C.4})$$

Since the entries of F are i.i.d. $\mathcal{N}(0, 1)$, we have large deviation inequality $|T_h^{-1}F'F' - 1| \lesssim_P T^{-1/2}$. This also implies that $\|F\| - T_h^{-1/2} \lesssim_P 1$ by Weyl's inequality.

Similarly, we can get $|T_h^{-1}\epsilon_1\epsilon_1' - 1| \lesssim_P T^{-1/2}$ and $\|\epsilon_1\| - T_h^{-1/2} \lesssim_P 1$. In addition, by Lemma A.1 in Wang and Fan (2017), we have $\|N^{-1}U'U - A_1'A_1\| \leq \|A_1\|^2 \|N^{-1}\epsilon'\epsilon - \mathbb{I}_{T_h}\| \lesssim_P \sqrt{T/N}$. Next, by direct calculation using the previous inequalities we obtain

$$\left\| \frac{F'\epsilon_1 A_1 + A_1'\epsilon_1' F}{T_h \sqrt{\lambda}} + \frac{U'U - N A_1' A_1}{T_h \lambda} \right\| \lesssim_P \frac{1}{\sqrt{\lambda}} + \frac{\sqrt{NT}}{T\lambda} \lesssim_P \frac{1}{\sqrt{\lambda}}.$$

Together with (C.3), we have

$$\left\| \frac{X'X}{T_h \lambda} - \frac{F'F}{T_h} - \frac{N A_1' A_1}{T_h \lambda} \right\| \lesssim_P \frac{1}{\sqrt{\lambda}}. \quad (\text{C.5})$$

Let η denote the first eigenvector of the matrix

$$M := T_h^{-1}F'F + \delta A_1' A_1.$$

With the assumption that $N/(T\lambda) \rightarrow \delta$, $(\lambda_1(M) - \lambda_2(M))/\lambda_1(M) \gtrsim_{\mathbf{P}} 1$ and (C.5), by the sin-theta theorem in Davis and Kahan (1970), we have $\|\mathbb{P}_\eta - \mathbb{P}_{\hat{\xi}}\| = \|\mathbb{P}_\eta - \mathbb{P}_{\hat{F}'}\| = o_{\mathbf{P}}(1)$.

In the case that $A_1' A_1 = \mathbb{I}_{T_h}$, the eigenvalues of M are given by

$$\lambda_i = \begin{cases} T_h^{-1} F F' + \delta & i = 1; \\ \delta & i \geq 2. \end{cases} \quad (\text{C.6})$$

and the first eigenvector is $F'/\|F\|$. Since the largest eigenvalue of $X'X/(T_h\lambda)$ is $\hat{\lambda}/\lambda$ with its corresponding eigenvector $\hat{\xi}$, (C.5) and Weyl's theorem yield that

$$\frac{\hat{\lambda}}{\lambda} = \frac{F F'}{T_h} + \frac{N}{T_h \lambda} + O_{\mathbf{P}}\left(\frac{1}{\sqrt{\lambda}}\right) = 1 + \delta + o_{\mathbf{P}}(1), \quad (\text{C.7})$$

and the sin-theta theorem implies that

$$\|\mathbb{P}_{F'} - \mathbb{P}_{\hat{\xi}}\| = \|F'(F F')^{-1} F - \hat{\xi} \hat{\xi}'\| = o_{\mathbf{P}}(1). \quad (\text{C.8})$$

Furthermore, (C.8) implies that $(F F')^{-1} (F \hat{\xi})^2 = \hat{\xi}' F' (F F')^{-1} F \hat{\xi} = 1 + o_{\mathbf{P}}(1)$. Together with $|T_h^{-1} F F' - 1| \lesssim T^{-1/2}$, and the fact that the sign of $\hat{\xi}$ plays no role in the estimator $\hat{y}_{T+h}$, we can choose $\hat{\xi}$ such that

$$\frac{F \hat{\xi}}{\sqrt{T_h}} - 1 = o_{\mathbf{P}}(1). \quad (\text{C.9})$$

Therefore, we have

$$\hat{y}_{t+h} = \hat{\alpha} \hat{\zeta}' x_T = \frac{Y \hat{\xi} \hat{\zeta}' x_T}{\sqrt{T_h \hat{\lambda}}} = \alpha \frac{F \hat{\xi} \hat{\zeta}' x_T}{\sqrt{T_h \hat{\lambda}}} = \alpha \frac{\hat{\zeta}' \beta f_T + \hat{\zeta}' u_T}{\sqrt{\hat{\lambda}}} (1 + o_{\mathbf{P}}(1)). \quad (\text{C.10})$$

Using (C.4), we have

$$\frac{\hat{\zeta}' \beta}{\sqrt{\hat{\lambda}}} = \frac{\sqrt{\lambda} \hat{\zeta}_1}{\sqrt{\hat{\lambda}}} = \frac{\lambda}{\hat{\lambda}} \frac{(F + \lambda^{-1/2} \epsilon_1 A_1) \hat{\xi}}{\sqrt{T_h}} = \frac{\lambda}{\hat{\lambda}} \left(\frac{F \hat{\xi}}{\sqrt{T_h}} + \frac{\epsilon_1 A_1 \hat{\xi}}{\sqrt{T_h \lambda}} \right).$$

Using (C.7), (C.9), $\|A_1\| \leq 1$, and $\|\epsilon_1\| \lesssim_{\mathbf{P}} \sqrt{T}$, it follows that

$$\frac{\hat{\zeta}' \beta}{\sqrt{\hat{\lambda}}} \xrightarrow{\mathbf{P}} \frac{1}{1 + \delta}. \quad (\text{C.11})$$

In addition, as $\text{Cov}(u_s, u_t) = 0$ for $s \neq t$, u_T is independent of $\hat{\zeta}$ and thus $\hat{\zeta}' u_T = O_{\mathbf{P}}(1)$. Combined

with (C.10) and (C.11), we have

$$\hat{y}_{T+h} \xrightarrow{P} \frac{\alpha f_T}{1+\delta} = (1+\delta)^{-1} E_T(y_{T+h}).$$

□

C.3 Proof of Propositions 3 and 4

Proof. In the case $d = K = 1$ and $z_t = 0$, the PLS estimate of the factor is $\hat{F} = FX'X$. With (C.5) and $T_h^{-1}FF' - 1 = o_P(T^{-1/2})$, we have

$$\left\| T_h^{-1} \lambda^{-1} \hat{F} - F (\mathbb{I}_{T_h} + \delta A_1' A_1) \right\| = o_P(T^{1/2}). \quad (\text{C.12})$$

Let $\eta = F (\mathbb{I}_{T_h} + \delta A_1' A_1)$, and $\hat{\xi}_1 = \hat{F} / \|\hat{F}\|$, $\hat{\xi}_2 = \eta / \|\eta\|$, with $\|\eta\| \asymp T^{1/2}$ and $\left\| T_h^{-1} \lambda^{-1} \hat{F} \right\| - \|\eta\| = o_P(T^{1/2})$ implied by (C.12), we have $\left\| \hat{\xi}_1 - \hat{\xi}_2 \right\| \xrightarrow{P} 0$ and thus

$$\left\| \mathbb{P}_{\hat{F}'} - \mathbb{P}_{\eta'} \right\| = \left\| \hat{\xi}_1' \hat{\xi}_1 - \hat{\xi}_2' \hat{\xi}_2 \right\| \leq 2 \left\| \hat{\xi}_1 - \hat{\xi}_2 \right\| \xrightarrow{P} 0.$$

This completes the proof of Proposition 3. In the special case $A_1' A_1 = \mathbb{I}_{T_h}$, as in Section 3.4.2, we can write

$$\hat{y}_{T+h} = \|YX'X\|^{-2} YX'XY'YX'x_T = \alpha \|FX'X\|^{-2} FX'XF'FX'x_T. \quad (\text{C.13})$$

We now analyze $\|FX'X\|$, $FX'XF'$, and $FX'x_T$, respectively. Recall that from (C.5), we have

$$\left\| \frac{X'X}{T_h \lambda} - \frac{F'F}{T_h} - \delta \mathbb{I}_{T_h} \right\| = o_P(1),$$

with $|T_h^{-1}FF' - 1| \lesssim_P T^{-1/2}$, we have

$$\frac{1}{T_h^{3/2} \lambda} \|FX'X\| = \frac{1}{\sqrt{T_h}} \left\| F \left(\frac{F'F}{T_h} + \delta \mathbb{I}_{T_h} \right) \right\| + o_P(1) = \frac{1}{\sqrt{T}} \left\| \frac{FF'F}{T} + \delta F \right\| + o_P(1) \xrightarrow{P} 1 + \delta. \quad (\text{C.14})$$

For the same reason, by direct calculation we have

$$\frac{1}{T_h^2 \lambda} FX'XF' = \frac{1}{T_h} F \left(\frac{F'F}{T_h} + \delta \mathbb{I}_{T_h} \right) F' + o_P(1) = \frac{FF'FF'}{T_h^2} + \delta \frac{FF'}{T_h} + o_P(1) \xrightarrow{P} 1 + \delta. \quad (\text{C.15})$$

Next, write X in the form of (C.3) as in the proof of Proposition 1. Then, using $\|\epsilon_1\| \lesssim_P \sqrt{T}$,

we have

$$\frac{1}{T_h \lambda} F X' \beta = \frac{F F'}{T_h} + \frac{F A_1' \epsilon_1'}{T \sqrt{\lambda}} \xrightarrow{P} 1. \quad (\text{C.16})$$

In addition, as u_T is independent of f_t and x_t for $t < T$, and (C.15), we have

$$\frac{1}{T_h \lambda} \|F X' u_T\| \lesssim_P \frac{1}{T_h \lambda} \|F X'\| \xrightarrow{P} 0. \quad (\text{C.17})$$

On the basis of (C.14), (C.15), (C.16), (C.17) and (C.13), we have concluded the proof. $\square$

C.4 Proof of Proposition 5

Proof. The explicit form of the PLS estimator in this case is

$$\hat{y}_{T+h}^{PLS} = \|Y X' X\|^{-2} Y X' X Y' Y X' x_T = \|Z U' U\|^{-2} Z U' U Z' Z U' u_T.$$

Recall that $U = (u_1, \dots, u_{T-h})$, u_T is independent of U and Z . Therefore,

$$\text{Var}(\hat{y}_{T+h}^{PLS}) = \|Z U' U\|^{-4} \|Z U'\|^6.$$

As z_i and u_i are generated from independent standard normal distribution, we have $\|Z U'\| \asymp_P T^{1/2} N^{1/2}$ and $\|U\| \asymp_P N^{1/2} + T^{1/2}$. Consequently,

$$\text{Var}(\hat{y}_{T+h}^{PLS}) \gtrsim_P \frac{N^3 T}{N^4 + T^4}.$$

On the other hand, the PCA estimator is

$$\hat{y}_{T+h}^{PCA} = \|U\|^{-1} Z \hat{\xi} \hat{\xi}' u_T,$$

where $\hat{\varsigma}$ and $\hat{\xi}$ are the first left and right singular vectors of U . Note that Z is independent of $\hat{\xi}$ and u_T is independent of $\hat{\varsigma}$, we have $\|Z \hat{\xi}\| \lesssim_P 1$ and $\|\hat{\varsigma}' u_T\| \lesssim_P 1$. Along with the fact that $\|U\| \gtrsim_P N^{1/2} + T^{1/2}$, we have $\hat{y}_{T+h}^{PCA} \lesssim_P 1/(N^{1/2} + T^{1/2})$. $\square$

C.5 Proof of Proposition 6

Proof. The estimated factor $\hat{F}$ is the first eigenvector of $X' X = U' U$. By Lemma A.1 in Wang and Fan (2017), we have $\|N^{-1} U' U - \text{diag}(1, \dots, 1, 1 + \epsilon)\| \lesssim_P \sqrt{T/N}$. Note that the first eigenvector of $\text{diag}(1, \dots, 1, 1 + \epsilon)$ is $(0, 0, \dots, 1)$, sin-theta theorem implies that $|\hat{f}_T| / \|\hat{F}\| \xrightarrow{P} 1$. As $\|\hat{F}\|^2 + \hat{f}_T^2 =$

$\|\widehat{F}\|^2$, we have $\|\widehat{F}\|/\widehat{f}_T \xrightarrow{P} 0$. As $\overline{Z}$ is independent of U , conditioning on U , the estimated coefficient

$$\widehat{\alpha} = \overline{Z}\widehat{F}' \left(\widehat{F}\widehat{F}' \right)^{-1}$$

follows a normal distribution with mean 0 and variance $\|\widehat{F}\|^{-2}$. Consequently,

$$\text{Var}(\widehat{f}_{T+h}^{SW}|U) = \text{Var}(\widehat{\alpha}\widehat{f}_T|U) = \left(\widehat{f}_T / \|\widehat{F}\| \right)^2 \xrightarrow{P} \infty,$$

which in turn implies that $\text{Var}(\widehat{f}_{T+h}^{SW}) \rightarrow \infty$. On the other hand, in our PCA algorithm, let $\widehat{\varsigma}$ and $\widehat{\xi}$ denote the first left and right singular vectors of $\underline{X} = \underline{U}$, then

$$\widehat{y}_{T+h}^{PCA} = \|\underline{U}\|^{-1} \overline{Z}\widehat{\xi}\widehat{\varsigma}' u_T.$$

Note that $\overline{Z}$ is independent of $\widehat{\xi}$ and u_T is independent of $\widehat{\varsigma}$, we have $\|\overline{Z}\widehat{\xi}\| \lesssim_P 1$ and $\|\widehat{\varsigma}' u_T\| \lesssim_P 1$. Along with the fact that $\|\underline{U}\| \gtrsim_P N^{1/2} + T^{1/2}$, we have $\widehat{y}_{T+h}^{PCA} \xrightarrow{P} 0$. $\square$

C.6 Technical Lemmas and Their Proofs

Without loss of generality, we assume that $\Sigma_f = \mathbb{I}_K$ in the following lemmas. Also, except for Lemma 3, we assume that $\widehat{K} = \widetilde{K}$ and $\widehat{I}_k = I_k$ for $k \leq \widetilde{K}$, which hold with probability approaching one as we will show in Lemma 3.

Lemma 1. *The singular vectors $\widehat{\xi}_{(k)}$ s we obtain from Algorithm 1 satisfy $\widehat{\xi}_{(j)}'\widehat{\xi}_{(k)} = \delta_{jk}$ for $j, k \leq \widehat{K}$.*

Proof. If $j = k$, this result holds from the definition of $\widehat{\xi}_{(k)}$. If $j < k$, recall that $\widetilde{X}_{(k)}$ is defined in (14) and $\widehat{\xi}_{(k)}$ is the first right singular vector of $\widetilde{X}_{(k)}$, we have

$$\widetilde{X}_{(k)} = X_{[I_k]} \prod_{i < k} \left(\mathbb{I}_T - \widehat{\xi}_{(i)}\widehat{\xi}_{(i)}' \right) \quad \text{and} \quad \widehat{\xi}_{(k)} = \arg \max_{v \in \mathbb{R}^T} \frac{\|\widetilde{X}_{(k)}v\|}{\|v\|}.$$

If $\widehat{\xi}_{(k)}'\widehat{\xi}_{(j)} = c_0 \neq 0$ for some $j < k$, then

$$\left\| \widetilde{X}_{(k)}(\widehat{\xi}_{(k)} - c_0\widehat{\xi}_{(j)}) \right\| = \left\| \widetilde{X}_{(k)}\widehat{\xi}_{(k)} - c_0\widetilde{X}_{(k)}\widehat{\xi}_{(j)} \right\| = \left\| \widetilde{X}_{(k)}\widehat{\xi}_{(k)} \right\|, \quad (\text{C.18})$$

since the definition of $\widetilde{X}_{(k)}$ implies that $\widetilde{X}_{(k)}\widehat{\xi}_{(j)} = 0$ for $j < k$.

On the other hand, since $\widehat{\xi}_{(k)}'\widehat{\xi}_{(j)} = c_0 \neq 0$, we have $(\widehat{\xi}_{(k)} - c_0\widehat{\xi}_{(j)})'\widehat{\xi}_{(j)} = 0$, and consequently,

$$\left\| \widehat{\xi}_{(k)} \right\|^2 = \left\| \widehat{\xi}_{(k)} - c_0\widehat{\xi}_{(j)} \right\|^2 + \left\| c_0\widehat{\xi}_{(j)} \right\|^2 > \left\| \widehat{\xi}_{(k)} - c_0\widehat{\xi}_{(j)} \right\|^2. \quad (\text{C.19})$$

Apparently, if $\|\tilde{X}_{(k)}\| = 0$, the SPCA procedure will terminate so we have $\|\tilde{X}_{(k)}\| > 0$ for $k \leq \hat{K}$. Together with (C.18) and (C.19), we have

$$\|\tilde{X}_{(k)}\| = \frac{\|\tilde{X}_{(k)}\hat{\xi}_{(k)}\|}{\|\hat{\xi}_{(k)}\|} \leq \frac{\|\tilde{X}_{(k)}(\hat{\xi}_{(k)} - c_0\hat{\xi}_{(j)})\|}{\|\hat{\xi}_{(k)} - c_0\hat{\xi}_{(j)}\|},$$

which contradicts with the definition of $\hat{\xi}_{(k)}$. Therefore, $\hat{\xi}_{(k)}'\hat{\xi}_{(j)} = 0$ for $j < k$. This completes the proof. $\square$

Lemma 2. *Under assumptions of Theorem 1, $b_{(k)}$, $\beta_{(k)}$ and $\tilde{K}$ defined in Section 3.1 satisfy*

$$(i) \quad b_{(j)}'b_{(k)} = \delta_{jk} \text{ for } j \leq k \leq \tilde{K}.$$

$$(ii) \quad \lambda_{(k)}^{1/2} = \|\beta_{(k)}\| \asymp q^{1/2}N^{1/2}.$$

$$(iii) \quad \tilde{K} \leq K.$$

$$(iv) \quad \tilde{K} = K, \text{ if we further have } \lambda_K(\alpha'\alpha) \gtrsim 1.$$

Proof. (i) Recall that $b_{(k)}$ is the first right singular vector of $\beta_{(k)}$ and $\beta_{(k)} = \beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}}$. Using the same argument as in the proof of Lemma 1, we have $b_{(j)}'b_{(k)} = \delta_{jk}$ for $j, k \leq \tilde{K}$.

(ii) The selection rule at k th step implies that

$$|I_k|^{-1} \sum_{i \in I_k} \left\| \beta_{[i]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}}^2 \geq N_0^{-1} \sum_{i \in I_0} \left\| \beta_{[i]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}}^2. \quad (\text{C.20})$$

For any matrix $A \in \mathbb{R}^{N \times D}$ and set $I \subset [N]$, we have

$$\sum_{i \in I} \|A_{[i]}\|_{\text{MAX}}^2 \leq \|A\|_{\text{F}}^2 \leq D \sum_{i \in I} \|A_{[i]}\|_{\text{MAX}}^2,$$

and $\|A\|^2 \leq \|A\|_{\text{F}}^2 \leq D \|A\|^2$. We thereby have

$$\|A\|^2 \asymp \sum_{i \in I} \|A_{[i]}\|_{\text{MAX}}^2. \quad (\text{C.21})$$

Using this result, (C.20) becomes

$$|I_k|^{-1} \left\| \beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|^2 \gtrsim N_0^{-1} \left\| \beta_{[I_0]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|^2.$$

Then, we have

$$\frac{\|\beta_{(k)}\|}{\sqrt{|I_k|}} \left\| \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\| \geq \frac{1}{\sqrt{|I_k|}} \left\| \beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\| \gtrsim \frac{1}{\sqrt{N_0}} \left\| \beta_{[I_0]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\| \geq \frac{\sigma_K(\beta_{[I_0]})}{\sqrt{N_0}} \left\| \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' \right\|, \quad (\text{C.22})$$

where we use $\beta_{[I_k]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha' = \beta_{[I_k]} (\prod_{j < k} \mathbb{M}_{b_{(j)}})^2 \alpha' = \beta_{(k)} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha'$ in the first inequality. With $\sigma_K(\beta_{[I_0]}) \gtrsim \sqrt{N_0}$ from Assumption 2, (C.22) leads to $\|\beta_{(k)}\| \gtrsim |I_k|^{1/2}$. In addition, $\|\beta\|_{\text{MAX}} \lesssim 1$ from Assumption 2 leads to $\|\beta_{(k)}\| \lesssim |I_k|^{1/2}$. Therefore, we have $\|\beta_{(k)}\| \asymp |I_k|^{1/2} \asymp q^{1/2} N^{1/2}$.

(iii) From (i), we have shown that $b_{(k)}$'s are pairwise orthogonal for $k \leq \tilde{K}$. It is impossible to have more than K pairwise orthogonal K dimensional vectors. Thus, $\tilde{K} \leq K$.

(iv) Recall that $\tilde{K}$ is defined in Section 3.1. Since the SPCA procedure stops at $\tilde{K} + 1$, we have at most $qN - 1$ rows of β satisfying $\left\| \beta_{[i]} \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}} \geq c$, which implies

$$\left\| \beta_{[I_0]} \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \alpha' \right\|^2 \lesssim qN + (N_0 - qN)c^2 = o(N_0),$$

where we use (C.21) and the assumptions $c \rightarrow 0$, $qN/N_0 \rightarrow 0$, and a similar argument for the proof of (20). With $\sigma_K(\beta_{[I_0]}) \gtrsim \sqrt{N_0}$ from Assumption 2, we have

$$\left\| \alpha \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \right\| \leq \sigma_K(\beta_{[I_0]})^{-1} \left\| \beta_{[I_0]} \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \alpha' \right\| = o(1). \quad (\text{C.23})$$

If $\tilde{K} \leq K - 1$, using (i), we have

$$\alpha \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} = \alpha - \alpha \sum_{j \leq \tilde{K}} b_{(j)} b'_{(j)},$$

which implies that

$$\sigma_K(\alpha) \leq \sigma_1 \left(\alpha \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \right) + \sigma_K \left(\alpha \sum_{j \leq \tilde{K}} b_{(j)} b'_{(j)} \right). \quad (\text{C.24})$$

Since

$$\text{Rank} \left(\alpha \sum_{j \leq \tilde{K}} b_{(j)} b'_{(j)} \right) \leq \tilde{K} \leq K - 1, \quad (\text{C.25})$$

we have $\sigma_K \left(\alpha \sum_{j \leq \tilde{K}} b_{(j)} b'_{(j)} \right) = 0$. Therefore, by (C.24) and (C.23), we further have $\sigma_K(\alpha) \lesssim$

$\sigma_1 \left(\alpha \prod_{j \leq \tilde{K}} \mathbb{M}_{b_{(j)}} \right) \rightarrow 0$. This contradicts with the assumption that $\lambda_K(\alpha' \alpha) \gtrsim 1$. Therefore, we have established that $\tilde{K} \geq K$. Together with the result in (iii), we have $\tilde{K} = K$. $\square$

Lemma 3. *Under assumptions of Theorem 1, for $k \leq \tilde{K}$, I_k , $\tilde{K}$ and $\beta_{(k)}$ defined in Section 3.1 satisfy*

$$(i) \ P(\hat{I}_k = I_k) \rightarrow 1.$$

$$(ii) \ \left\| \tilde{X}_{(k)} - \beta_{(k)} F \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}.$$

$$(iii) \ |\hat{\lambda}_{(k)}^{1/2} / \lambda_{(k)}^{1/2} - 1| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}, \text{ and } \hat{\lambda}_{(k)} \asymp_P \lambda_{(k)} \asymp qN.$$

$$(iv) \ \left\| T_h^{-1/2} F \hat{\xi}_{(k)} - b_{(k)} \right\| \asymp \left\| \mathbb{P}_{\hat{F}'_{(k)}} - T_h^{-1} F' \mathbb{P}_{b_{(k)}} F \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}.$$

$$(v) \ P(\hat{K} = \tilde{K}) \rightarrow 1.$$

For $k \leq \tilde{K} + 1$, we have

$$(vi) \ \left\| T_h^{-1} X \prod_{j=1}^{k-1} Y'_{(j)} - \beta \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}} \lesssim_P (\log NT)^{1/2} (q^{-1/2} N^{-1/2} + T^{-1/2}).$$

Proof. We prove (i)-(iv) by induction. First, we show that (i)-(iv) hold when $k = 1$:

(i) Recall that $\hat{I}_1$ is selected based on $T_h^{-1} X Y'$ and I_1 is selected based on $\beta \alpha'$. With simple algebra, we have

$$T_h^{-1} X Y' - \beta \alpha' = \beta (T_h^{-1} F F' - \mathbb{I}_K) \alpha' + T_h^{-1} U F' \alpha' + T_h^{-1} \beta F Z' + T_h^{-1} U Z'.$$

With Assumptions 1, 2 and 4, we have

$$\begin{aligned} \|T_h^{-1} X Y' - \beta \alpha'\|_{\text{MAX}} &\lesssim \|\beta\|_{\text{MAX}} \|T_h^{-1} F F' - \mathbb{I}_K\| \|\alpha'\| + T_h^{-1} \|U F'\|_{\text{MAX}} \|\alpha'\| \\ &\quad + T_h^{-1} \|\beta\|_{\text{MAX}} \|F Z'\| + T_h^{-1} \|U Z'\|_{\text{MAX}} \lesssim_P (\log N)^{1/2} T^{-1/2}. \end{aligned}$$

From Assumption 5, we have $c_{qN}^{(1)} - c_{qN+1}^{(1)} \gtrsim c_{qN}^{(1)}$ and the definition of $\tilde{K}$ implies that $c_{qN}^{(k)} \geq c$ for $k \leq \tilde{K}$. Thus, we have $c_{qN}^{(1)} - c_{qN+1}^{(1)} \gtrsim c$. Define the events

$$\begin{aligned} A_1 &:= \left\{ \|T_h^{-1} X_{[i]} Y'\|_{\text{MAX}} > (c_{qN}^{(1)} + c_{qN+1}^{(1)})/2 \text{ for all } i \in I_1 \right\}, \\ A_2 &:= \left\{ \|T_h^{-1} X_{[i]} Y'\|_{\text{MAX}} < (c_{qN}^{(1)} + c_{qN+1}^{(1)})/2 \text{ for all } i \in I_1^c \right\}, \\ A_3 &:= \left\{ \|T_h^{-1} X_{[i]} Y' - \beta_{[i]} \alpha'\|_{\text{MAX}} \geq (c_{qN}^{(1)} - c_{qN+1}^{(1)})/2 \text{ for some } i \in [N] \right\}. \end{aligned} \quad (\text{C.26})$$

It is easy to observe that $\{\hat{I}_1 = I_1\} \supset A_1 \cap A_2$. In addition, from the definition of I_1 , we have $\|\beta_{[i]} \alpha'\|_{\text{MAX}} \geq c_{qN}^{(1)}$ for all $i \in I_1$ and $\|\beta_{[i]} \alpha'\|_{\text{MAX}} \leq c_{qN+1}^{(1)}$ for all $i \in I_1^c$. Therefore, if A_1^c occurs, we have

$$\|T_h^{-1} X_{[i]} Y' - \beta_{[i]} \alpha'\|_{\text{MAX}} \geq (c_{qN}^{(1)} - c_{qN+1}^{(1)})/2,$$

for some $i \in I_1$, which implies $A_1^c \subset A_3$. Similarly, we have $A_2^c \subset A_3$. Using $\{\widehat{I}_1 = I_1\} \supset A_1 \cap A_2$ and $A_1^c \cup A_2^c \subset A_3$, we have

$$\mathbb{P}(\widehat{I}_1 = I_1) \geq \mathbb{P}(A_1 \cap A_2) = 1 - \mathbb{P}(A_1^c \cup A_2^c) \geq 1 - \mathbb{P}(A_3). \quad (\text{C.27})$$

Using $c^{-1}(\log N)^{1/2}T^{-1/2} \rightarrow 0$ and $c_{qN}^{(1)} - c_{qN+1}^{(1)} \gtrsim c$, we have $\mathbb{P}(A_3) \rightarrow 0$ and consequently, $\mathbb{P}(\widehat{I}_1 = I_1) \rightarrow 1$.

(ii) Since $\widehat{I}_1 = I_1$ with probability approaching one, we impose $\widehat{I}_1 = I_1$ below. Then, we have $\widetilde{X}_{(1)} = X_{[I_1]}$ by (14) and Assumption 3 gives $\|\widetilde{X}_{(1)} - \beta_{(1)}F\| = \|U_{[I_1]}\| \lesssim_P q^{1/2}N^{1/2} + T^{1/2}$.

(iii) From Lemma 13, we have $\sigma_j(\beta_{(1)}F)/\sigma_j(\beta_{(1)}) = T_h^{1/2} + O_P(1)$, which leads to

$$|\|\beta_{(1)}F\| - T_h^{1/2}\lambda_{(1)}^{1/2}| = |\sigma_1(\beta_{(1)}F) - T_h^{1/2}\sigma_1(\beta_{(1)})| \lesssim_P q^{1/2}N^{1/2}, \quad (\text{C.28})$$

where we use $\lambda_{(1)}^{1/2} = \|\beta_{(1)}\| \asymp q^{1/2}N^{1/2}$ from Lemma 2 in the last step. In addition, the result in (ii) implies that

$$|\|\widetilde{X}_{(1)}\| - \|\beta_{(1)}F\|| \leq \|\widetilde{X}_{(1)} - \beta_{(1)}F\| \lesssim_P q^{1/2}N^{1/2} + T^{1/2}. \quad (\text{C.29})$$

Using (C.28), (C.29) and $\lambda_{(1)}^{1/2} \asymp q^{1/2}N^{1/2}$, we have

$$|\frac{\widehat{\lambda}_{(1)}^{1/2}}{\lambda_{(1)}^{1/2}} - 1| = |\frac{\|\widetilde{X}_{(1)}\|}{T_h^{1/2}\lambda_{(1)}^{1/2}} - 1| \leq \frac{|\|\widetilde{X}_{(1)}\| - \|\beta_{(1)}F\||}{T_h^{1/2}\lambda_{(1)}^{1/2}} + \frac{|\|\beta_{(1)}F\| - T_h^{1/2}\lambda_{(1)}^{1/2}|}{T_h^{1/2}\lambda_{(1)}^{1/2}} \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}.$$

and thus $\widehat{\lambda}_{(1)} \asymp_P qN$.

(iv) Let $\tilde{\xi}_{(1)} \in \mathbb{R}^{T_h \times 1}$ denote the first right singular vector of $\beta_{(1)}F$. From Lemma 13, we have

$$\|\mathbb{P}_{\tilde{\xi}_{(1)}} - T_h^{-1}F'\mathbb{P}_{b_{(1)}}F\| \lesssim_P T^{-1/2} \quad (\text{C.30})$$

and $\sigma_j(\beta_{(1)}F)/\sigma_j(\beta_{(1)}) = T_h^{1/2} + O_P(1)$. The latter further leads to

$$\sigma_1(\beta_{(1)}F) - \sigma_2(\beta_{(1)}F) = T_h^{1/2}(\sigma_1(\beta_{(1)}) - \sigma_2(\beta_{(1)})) + O_P(\sigma_1(\beta_{(1)})) \asymp_p T^{1/2}\sigma_1(\beta_{(1)}), \quad (\text{C.31})$$

where we use the assumption that $\sigma_2(\beta_{(1)}) \leq (1 + \delta)^{-1}\sigma_1(\beta_{(1)})$ in the last equation.

Using $\|\widetilde{X}_{(1)} - \beta_{(1)}F\| \lesssim_P q^{1/2}N^{1/2} + T^{1/2}$ as proved in (ii), (C.31), Lemma 2 and Wedin (1972)'s sin-theta theorem for singular vectors, we have

$$\|\mathbb{P}_{\widehat{F}_{(1)}} - \mathbb{P}_{\tilde{\xi}_{(1)}}\| \lesssim_P \frac{q^{1/2}N^{1/2} + T^{1/2}}{\sigma_1(\beta_{(1)}F) - \sigma_2(\beta_{(1)}F)} \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}. \quad (\text{C.32})$$

In light of (C.30) and (C.32), we have the first equation in (iv) holds for $k = 1$. As $\mathbb{P}_{\widehat{F}'_{(k)}} = \widehat{\xi}_{(k)}\widehat{\xi}'_{(k)}$, left and right multiplying this equation by $\widehat{\xi}'_{(1)}$ and $\widehat{\xi}_{(1)}$, we have

$$|1 - T_h^{-1}(b'_{(1)}F\widehat{\xi}_{(1)})^2| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2},$$

which leads to $|1 - T_h^{-1/2}b'_{(1)}F\widehat{\xi}_{(1)}| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}$. Left-multiplying it by $b_{(1)}$ gives the second equation in (iv).

So far, we have proved that (i)-(iv) hold for $k = 1$. Now, assuming that (i)-(iv) hold for $j \leq k-1$, we will show that (i)-(iv) continue to hold for $j = k$.

(i) Again, we show the difference between the sample covariances and their population counterparts introduced in the SPCA procedure is tiny. At the k th step, the difference can be written as

$$\begin{aligned} & \left\| \beta \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \alpha' - T_h^{-1}(\beta F + U) \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} (\alpha F + Z)' \right\|_{\text{MAX}} \\ & \leq \left\| \beta \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \alpha' - T_h^{-1} \beta F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' \alpha' \right\|_{\text{MAX}} + T_h^{-1} \left\| \beta F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} Z' \right\|_{\text{MAX}} \\ & \quad + T_h^{-1} \left\| U \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' \alpha' \right\|_{\text{MAX}} + T_h^{-1} \left\| U \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} Z' \right\|_{\text{MAX}}. \end{aligned} \quad (\text{C.33})$$

Since (iv) holds for $j \leq k-1$, we have

$$\left\| \sum_{j=1}^{k-1} \mathbb{P}_{\widehat{F}'_{(j)}} - T_h^{-1} F' \sum_{j=1}^{k-1} \mathbb{P}_{b_{(j)}} F \right\| = \left\| \sum_{j=1}^{k-1} \left(\mathbb{P}_{\widehat{F}'_{(j)}} - T_h^{-1} F' \mathbb{P}_{b_{(j)}} F \right) \right\| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}. \quad (\text{C.34})$$

Using Lemma 1 and Lemma 2(i), we have

$$\prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} = \mathbb{I}_K - \sum_{j=1}^{k-1} \mathbb{P}_{b_{(j)}}, \quad \text{and} \quad \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} = \mathbb{I}_{T_h} - \sum_{j=1}^{k-1} \mathbb{P}_{\widehat{F}'_{(j)}}.$$

Using the above equations, (C.34), and $\|T_h^{-1}FF' - \mathbb{I}_K\| \lesssim_P T^{-1/2}$, we have

$$T_h^{-1/2} \left\| F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} F \right\| = T_h^{-1/2} \left\| F \sum_{j=1}^{k-1} \mathbb{P}_{\widehat{F}'_{(j)}} - \sum_{j=1}^{k-1} \mathbb{P}_{b_{(j)}} F \right\| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}. \quad (\text{C.35})$$

Similarly, right multiplying F' to the term inside the $\|\cdot\|$ of (C.35), we have

$$\left\| T_h^{-1} F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}. \quad (\text{C.36})$$

Next, we analyze the four terms in (C.33) one by one. For the first term, using (C.36) and Assumption 2, we have

$$\begin{aligned} \left\| \beta \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \alpha' - T_h^{-1} \beta F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' \alpha' \right\|_{\text{MAX}} &\lesssim \|\beta\|_{\text{MAX}} \left\| \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} - T_h^{-1} F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' \right\| \|\alpha\| \\ &\lesssim_p q^{-1/2} N^{-1/2} + T^{-1/2}. \end{aligned}$$

For the second term, using (C.35), Assumption 2, we have

$$\begin{aligned} &T_h^{-1} \left\| \beta F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} Z' \right\|_{\text{MAX}} \\ &\lesssim T_h^{-1} \|\beta\|_{\text{MAX}} \left\| \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \right\| \|F Z'\| + T_h^{-1} \|\beta\|_{\text{MAX}} \left\| F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} F \right\| \|Z\| \\ &\lesssim_p q^{-1/2} N^{-1/2} + T^{-1/2}. \end{aligned}$$

For the third term, using (C.35), we have

$$\begin{aligned} &T_h^{-1} \left\| U \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} F' \alpha' \right\|_{\text{MAX}} \\ &\lesssim T_h^{-1} \|U F'\|_{\text{MAX}} \left\| \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \right\| \|\alpha\| + T_h^{-1} \|U\|_{\text{MAX}} T_h^{1/2} \left\| F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} F \right\| \|\alpha\| \\ &\lesssim_p (\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right). \end{aligned}$$

For the forth term, using (C.34), we have

$$\begin{aligned} &T_h^{-1} \left\| U \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} Z' \right\|_{\text{MAX}} \lesssim T_h^{-1} \|U Z'\|_{\text{MAX}} + T_h^{-2} \|U F'\|_{\text{MAX}} \left\| \sum_{j=1}^{k-1} \mathbb{P}_{b_{(j)}} \right\| \|F Z'\| \\ &\quad + T_h^{-1/2} \|U\|_{\text{MAX}} \left\| T_h^{-1} F' \sum_{j=1}^{k-1} \mathbb{P}_{b_{(j)}} F - \sum_{j=1}^{k-1} \mathbb{P}_{\widehat{F}'_{(j)}} \right\| \|Z\| \\ &\lesssim_p (\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right). \end{aligned}$$

Hence, we have

$$\left\| T_h^{-1} X \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} Y' - \beta \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}} \lesssim_P (\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right). \quad (\text{C.37})$$

As in the case of $k = 1$, with the assumption that

$$c^{-1}(\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right) \rightarrow 0,$$

and Assumption 5, we can reuse the arguments for (C.26) and (C.27) in the case of $k = 1$ and obtain $P(\widehat{I}_k = I_k) \rightarrow 1$.

(ii) We impose $\widehat{I}_k = I_k$ below. Then, we have $\widetilde{X}_{(k)} = X_{[I_k]} \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}}$ and thus

$$\widetilde{X}_{(k)} - \beta_{(k)} F = X_{[I_k]} \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \beta_{(k)} F = \beta_{[I_k]} \left(F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} F \right) + U_{[I_k]} \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}}.$$

Hence, using Assumption 2 and (C.35), we have

$$\left\| \widetilde{X}_{(k)} - \beta_{(k)} F \right\| \leq \|\beta_{[I_k]}\| \left\| F \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} - \prod_{j=1}^{k-1} \mathbb{M}_{b_{(j)}} F \right\| + \|U_{[I_k]}\| \left\| \prod_{j=1}^{k-1} \mathbb{M}_{\widehat{F}'_{(j)}} \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}.$$

(iii) The proof of (iii) is analogous to the case $k = 1$.

(iv) The proof of (iv) is analogous to the case $k = 1$.

To sum up, by induction, we have shown that (i)-(iv) hold for $k \leq \widetilde{K}$.

(v) Recall that $\widetilde{K}$ is determined by $\beta_{[i]} \prod_{j < k} \mathbb{M}_{b_{(j)}} \alpha'$ whereas $\widehat{K}$ is determined by $T_h^{-1} X_{[i]} \prod_{j < k} \mathbb{M}_{\widehat{F}'_{(j)}} Y'$. Since (iv) holds for $j \leq \widetilde{K}$ as shown above, using the same proof for (C.37), we have

$$\left\| T_h^{-1} X \prod_{j=1}^{\widetilde{K}} \mathbb{M}_{\widehat{F}'_{(j)}} Y' - \beta \prod_{j=1}^{\widetilde{K}} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}} \lesssim_P (\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right). \quad (\text{C.38})$$

The assumption $c_{qN}^{(\widetilde{K}+1)} \leq (1 + \delta)^{-1} c$ in Assumption 5 implies that $c - c_{qN}^{(\widetilde{K}+1)} \asymp c$. Together with

$$c^{-1}(\log NT)^{1/2} \left(q^{-1/2} N^{-1/2} + T^{-1/2} \right) \rightarrow 0,$$

we can reuse the arguments for (C.26) and (C.27) with events

$$B_1 := \left\{ \left\| T_h^{-1} X_{[i]} \prod_{j=1}^{\widetilde{K}} \mathbb{M}_{\widehat{F}'_{(j)}} Y' \right\|_{\text{MAX}} > (c + c_{qN}^{(\widetilde{K}+1)})/2 \text{ for at most } qN - 1 \text{ different } i \text{ in } [N] \right\},$$

$$B_2 := \left\{ \left\| T_h^{-1} X_{[i]} \prod_{j=1}^{\tilde{K}} \mathbb{M}_{\hat{F}_{(j)}} Y' - \beta_{[i]} \prod_{j=1}^{\tilde{K}} \mathbb{M}_{b_{(j)}} \alpha' \right\|_{\text{MAX}} \geq (c - c_{qN}^{(\tilde{K}+1)})/2 \text{ for some } i \in [N] \right\}, \quad (\text{C.39})$$

to obtain $\mathbb{P}(\hat{K} = \tilde{K}) \geq \mathbb{P}(B_1) = 1 - \mathbb{P}(B_1^c) \geq 1 - \mathbb{P}(B_2) \rightarrow 1$.

(vi) This result comes directly from (C.37) and (C.38). $\square$

Lemma 4. *Under assumptions of Theorem 1, for $k \leq \tilde{K}$, we have*

- (i) $\left\| U'_{[I_k]} \hat{\varsigma}_{(k)} \right\| \lesssim_P T^{1/2} + T^{-1/2} q^{1/2} N^{1/2}.$
- (ii) $\left\| F U'_{[I_k]} \hat{\varsigma}_{(k)} \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}, \quad \left\| Z U'_{[I_k]} \hat{\varsigma}_{(k)} \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}.$
- (iii) $|\hat{\varsigma}'_{(k)}(u_T)_{[I_k]}| \lesssim_P 1 + T^{-1/2} q^{1/2} N^{1/2}.$

Proof. (i) Using Lemma 1, we have

$$T_h^{1/2} \hat{\lambda}_{(k)}^{1/2} \hat{\varsigma}_{(k)} = X_{[I_k]} \prod_{j=1}^{k-1} \mathbb{M}_{\hat{\xi}_{(j)}} \hat{\xi}_{(k)} = X_{[I_k]} \hat{\xi}_{(k)} = \beta_{[I_k]} F \hat{\xi}_{(k)} + U_{[I_k]} \hat{\xi}_{(k)}. \quad (\text{C.40})$$

Therefore, along with Assumption 1, Assumption 3 and Assumption 4(ii), we obtain

$$\begin{aligned} T_h^{1/2} \hat{\lambda}_{(k)}^{1/2} \left\| \hat{\varsigma}_{(k)} U_{[I_k]} \right\| &\leq \left\| \hat{\xi}_{(k)} F' \beta'_{[I_k]} U_{[I_k]} \right\| + \left\| \hat{\xi}_{(k)} U'_{[I_k]} U_{[I_k]} \right\| \\ &\leq \|F\| \left\| \beta'_{[I_k]} U_{[I_k]} \right\| + \left\| U'_{[I_k]} U_{[I_k]} \right\| \lesssim_P q^{1/2} N^{1/2} T + qN. \end{aligned} \quad (\text{C.41})$$

Together with $\hat{\lambda}_{(k)} \asymp_P qN$, we have the desired result.

(ii) Similarly, by Assumption 1, Assumption 3 and Assumption 4(i)(ii), we have

$$\begin{aligned} T^{1/2} \hat{\lambda}_{(k)}^{1/2} \left\| \hat{\varsigma}_{(k)} U_{[I_k]} F' \right\| &\leq \left\| \hat{\xi}_{(k)} F' \beta'_{[I_k]} U_{[I_k]} F' \right\| + \left\| \hat{\xi}_{(k)} U'_{[I_k]} U_{[I_k]} F' \right\| \\ &\leq \|F\| \left\| \beta'_{[I_k]} U_{[I_k]} F' \right\| + \|U_{[I_k]}\| \left\| U_{[I_k]} F' \right\| \lesssim_P q^{1/2} N^{1/2} T + qNT^{1/2}. \end{aligned} \quad (\text{C.42})$$

Together with $\hat{\lambda}_{(k)} \asymp_P qN$, we have the desired result. In addition, replacing F by Z and W , we have the second and third equations in (ii).

(iii) By Assumption 1, Assumption 3 and Assumption 4(iii), we have

$$\begin{aligned} T_h^{1/2} \hat{\lambda}_{(k)}^{1/2} |\hat{\varsigma}'_{(k)}(u_T)_{[I_k]}| &\leq |\hat{\xi}_{(k)} F' \beta'_{[I_k]}(u_T)_{[I_k]}| + |\hat{\xi}_{(k)} U'_{[I_k]}(u_T)_{[I_k]}| \\ &\leq \|F\| \left\| \beta'_{[I_k]}(u_T)_{[I_k]} \right\| + \|U_{[I_k]}\| \left\| (u_T)_{[I_k]} \right\| \lesssim_P q^{1/2} N^{1/2} T^{1/2} + qN. \end{aligned}$$

Together with $\hat{\lambda}_{(k)} \asymp_P qN$, we have the desired result. $\square$

Lemma 5. *Under assumptions of Theorem 1, for $k, l \leq \tilde{K}$, we have*

$$\begin{aligned}
(i) \quad & \left\| \frac{\tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}, \quad \left\| \frac{\tilde{U}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}. \\
(ii) \quad & \left\| \frac{F \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}, \quad \left\| \frac{Z \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}, \quad \left\| \frac{W \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}. \\
(iii) \quad & \left| \frac{\hat{\xi}_{(l)} U'_{[I_k]} \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right| \lesssim_P q^{-1} N^{-1} + T^{-1}, \quad \left| \frac{\hat{\xi}_{(l)} \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right| \lesssim_P q^{-1} N^{-1} + T^{-1}. \\
(iv) \quad & |\zeta'_{(k)} \tilde{D}_{(k)} u_T| \lesssim_P 1 + T^{-1/2} q^{1/2} N^{1/2}, \quad |\zeta'_{(k)} D_{(k)} u_T| \lesssim_P 1.
\end{aligned}$$

Proof. (i) Recall that from the definition of $U_{(k)}$ (below (16)), we have

$$\tilde{U}_{(k)} = U_{[I_k]} - \sum_{i=1}^{k-1} \frac{X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h}} \frac{\zeta'_{(i)} \tilde{U}_{(i)}}{\sqrt{\hat{\lambda}_{(i)}}}. \quad (\text{C.43})$$

Then, a direct multiplication of $\zeta'_{(k)}/\sqrt{T_h \hat{\lambda}_{(k)}}$ from the left side of (C.43) leads to

$$\frac{\zeta'_{(k)} \tilde{U}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} = \frac{\zeta'_{(k)} U_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} - \sum_{i=1}^{k-1} \frac{\zeta'_{(k)} X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \frac{\zeta'_{(i)} \tilde{U}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}}.$$

Consequently, using $\|X_{[I_k]}\| \leq \|\beta_{[I_k]}\| \|F\| + \|U_{[I_k]}\| \lesssim_P q^{1/2} N^{1/2} T^{1/2}$, $\hat{\lambda}_{(k)} \asymp_P qN$ and Lemma 4(i) we have

$$\left\| \frac{\zeta'_{(k)} \tilde{U}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \leq \left\| \frac{\zeta'_{(k)} U_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| + \sum_{i=1}^{k-1} \left\| \frac{X_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \left\| \frac{\zeta'_{(i)} \tilde{U}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1} + \sum_{i=1}^{k-1} \left\| \frac{\zeta'_{(i)} \tilde{U}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\|. \quad (\text{C.44})$$

If $\left\| T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} \zeta'_{(i)} \tilde{U}_{(i)} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}$ holds for $i \leq k-1$, then (C.44) implies that this inequality also holds for k . In addition, when $k=1$, $\tilde{U}_{(1)} = U_{[I_1]}$ and this equation is implied from Lemma 4(i). Therefore, we have (i) holds for $k \leq \tilde{K}$ by induction.

Using (C.43) again, with Assumption 3, we have

$$\left\| \frac{\tilde{U}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \leq \left\| \frac{U_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| + \sum_{i=1}^{k-1} \left\| \frac{X_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \left\| \frac{\tilde{U}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2} + \sum_{i=1}^{k-1} \left\| \frac{\tilde{U}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\|. \quad (\text{C.45})$$

When $k=1$, Assumption 3 implies that $\left\| T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \tilde{U}_{(k)} \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1/2}$. Then, using the same induction argument with (C.45), we have this inequality holds for $k \leq \tilde{K}$.

(ii) Similarly, by simple multiplication of F' from the right side of (C.43), we have

$$\frac{\zeta'_{(k)} \tilde{U}_{(k)} F'}{T_h \sqrt{\hat{\lambda}_{(k)}}} = \frac{\zeta'_{(k)} U_{[I_k]} F'}{T_h \sqrt{\hat{\lambda}_{(k)}}} - \sum_{i=1}^{k-1} \frac{\zeta'_{(k)} X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \frac{\zeta'_{(i)} \tilde{U}_{(i)} F'}{T_h \sqrt{\hat{\lambda}_{(i)}}}.$$

Consequently, we have

$$\begin{aligned} \left\| \frac{\zeta'_{(k)} \tilde{U}_{(k)} F'}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| &\leq \left\| \frac{\zeta'_{(k)} U_{[I_k]} F'}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| + \sum_{i=1}^{k-1} \left\| \frac{X_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \left\| \frac{\zeta'_{(i)} \tilde{U}_{(i)} F'}{T_h \sqrt{\hat{\lambda}_{(i)}}} \right\| \\ &\lesssim_P q^{-1} N^{-1} + T^{-1} + \sum_{i=1}^{k-1} \left\| \frac{\zeta'_{(i)} \tilde{U}_{(i)} F'}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\|. \end{aligned} \quad (\text{C.46})$$

When $k = 1$, $\left\| T_h^{-1} \hat{\lambda}_{(k)}^{-1/2} \zeta'_{(k)} \tilde{U}_{(k)} F' \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}$ is a result of Lemma 4(ii). Then, a direct induction argument using (C.46) leads to this inequality for $k \leq \tilde{K}$.

Replacing F by Z in the above proof, and using Lemma 4(ii), we have the following result:

$$\left\| \frac{Z \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{T \sqrt{\hat{\lambda}_{(k)}}} \right\| \lesssim_P q^{-1} N^{-1} + T^{-1}.$$

(iii) Recall that $\tilde{X}_{(k)} = \tilde{\beta}_{(k)} F + \tilde{U}_{(k)}$ as defined in (14), we have

$$|\zeta'_{(l)} \tilde{X}_{(l)} U'_{[I_k]} \hat{\varsigma}_{(k)}| \leq |\zeta'_{(l)} \tilde{\beta}_{(l)} F U'_{[I_k]} \hat{\varsigma}_{(k)}| + |\zeta'_{(l)} \tilde{U}_{(l)} U'_{[I_k]} \hat{\varsigma}_{(k)}| \leq \left\| \zeta'_{(l)} \tilde{\beta}_{(l)} \right\| \left\| F U'_{[I_k]} \hat{\varsigma}_{(k)} \right\| + \left\| \zeta'_{(l)} \tilde{U}_{(l)} \right\| \left\| U'_{[I_k]} \hat{\varsigma}_{(k)} \right\|.$$

Along with (12), we have

$$\left| \frac{\hat{\xi}'_{(l)} U'_{[I_k]} \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right| = \left| \frac{\zeta'_{(l)} \tilde{X}_{(l)} U'_{[I_k]} \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right| \leq \left\| \frac{\zeta'_{(l)} \tilde{\beta}_{(l)}}{\sqrt{\hat{\lambda}_{(l)}}} \right\| \left\| \frac{F U'_{[I_k]} \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}} \right\| + \left\| \frac{U'_{[I_k]} \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} \right\| \left\| \frac{\tilde{U}'_{(l)} \hat{\varsigma}_{(l)}}{\sqrt{T_h \hat{\lambda}_{(l)}}} \right\|. \quad (\text{C.47})$$

Using $\left\| \hat{\lambda}_{(k)}^{-1/2} \zeta'_{(k)} \tilde{\beta}_{(k)} \right\| \lesssim_P 1$ from Lemma 10, results of (i)(ii) and Lemma 4(i) completes the proof.

Replacing $U_{[I_k]}$ by $\tilde{U}_{(k)}$ above and using the inequality that

$$|\zeta'_{(l)} \tilde{X}_{(l)} \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}| \leq |\zeta'_{(l)} \tilde{\beta}_{(l)} F \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}| + |\zeta'_{(l)} \tilde{U}_{(l)} \tilde{U}'_{(k)} \hat{\varsigma}_{(k)}| \leq \left\| \zeta'_{(l)} \tilde{\beta}_{(l)} \right\| \left\| F \tilde{U}'_{(k)} \hat{\varsigma}_{(k)} \right\| + \left\| \zeta'_{(l)} \tilde{U}_{(l)} \right\| \left\| \tilde{U}'_{(k)} \hat{\varsigma}_{(k)} \right\|$$

and (12), we obtain the second equation in (iii).

(iv) Similar to (ii), by induction, we have

$$|\zeta'_{(k)} \tilde{D}_{(k)} u_T| \leq |\hat{\varsigma}_{(k)}(u_T)_{[I_k]}| + \sum_{i=1}^{k-1} \left\| \frac{X_{[I_k]}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| |\zeta'_{(i)} \tilde{D}_{(i)} u_T| \lesssim_P 1 + T^{-1/2} q^{1/2} N^{1/2}.$$

For the second inequality, from Lemma 2, we have $b'_{(i)}b_{(j)} = 0$ when $i \neq j$. Thus, by definition, we have $\varsigma_{(k)} = \lambda_{(k)}^{-1/2} \beta_{(k)} b_{(k)} = \lambda_{(k)}^{-1/2} \beta_{[I_k]} b_{(k)}$. Using $\|\beta'_{[I_k]}(u_T)_{[I_k]}\| \lesssim_P q^{1/2} N^{1/2}$ from Assumption 4(iii), $\lambda_{(k)} \asymp qN$ from Lemma 2 and the definition of $D_{(k)}$, we have

$$\begin{aligned} |\varsigma'_{(k)} D_{(k)} u_T| &\leq |\varsigma'_{(k)}(u_T)_{[I_k]}| + \sum_{i < k} \lambda_{(i)}^{-1/2} \|\beta_{[I_k]}\| \|b_{(i)}\| |\varsigma'_{(i)} D_{(i)} u_T| \\ &\leq \lambda_{(k)}^{-1/2} \|b_{(k)}\| \|\beta'_{[I_k]}(u_T)_{[I_k]}\| + \sum_{i < k} \lambda_{(i)}^{-1/2} \|\beta_{[I_k]}\| |\varsigma'_{(i)} D_{(i)} u_T| \\ &\lesssim_P 1 + \sum_{i < k} |\varsigma'_{(i)} D_{(i)} u_T|. \end{aligned}$$

As $|\varsigma'_{(1)} D_{(1)} u_T| \leq \lambda_{(1)}^{-1/2} \|b_{(1)}\| \|\beta_{[I_1]}(u_T)_{[I_1]}\| \lesssim_P 1$, we have $|\varsigma'_{(k)} D_{(k)} u_T| \lesssim_P 1$ by induction. $\square$

Lemma 6. *Under assumptions of Theorem 1, for $k \leq \tilde{K}$, we have*

$$\begin{aligned} (i) \quad &\|\hat{\xi}_{(k)} - T_h^{-1/2} F' b_{k2}\| \lesssim_P T^{-1} + q^{-1/2} N^{-1/2}. \\ (ii) \quad &\|T_h^{-1/2} Z \hat{\xi}_{(k)} - T_h^{-1} Z F' b_{k2}\| \lesssim_P T^{-1} + q^{-1} N^{-1}. \end{aligned}$$

Proof. (i) By the definitions of b_{k2} and $\hat{\xi}_{(k)}$ by (31) and (12), $\tilde{X}_{(k)} = \tilde{\beta}_{(k)} F + \tilde{U}_{(k)}$, we have

$$\hat{\xi}_{(k)} - T_h^{-1/2} F' b_{k2} = \frac{\tilde{U}'_{(k)} \hat{\varsigma}_{(k)}}{\sqrt{T \hat{\lambda}_{(k)}}}. \quad (\text{C.48})$$

Then, Lemma 5(i) leads to (i) directly.

(ii) Similarly, Lemma 5(ii) yields (ii). $\square$

Lemma 7. *Under assumptions of Theorem 1, for $k, j \leq \tilde{K}$, we have*

$$\begin{aligned} (i) \quad &\|\beta'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)}\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}. \\ (ii) \quad &\|\beta'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(j)}\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}. \\ (iii) \quad &\|(u_T)'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)}\| \lesssim_P T^{-1/2} q N + T^{1/2}. \\ (iv) \quad &\|(u_T)'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(j)}\| \lesssim_P T^{-1/2} q N + T^{1/2}. \end{aligned}$$

Proof. (i) With Lemma 6 and Assumption 4, we have

$$\|\beta'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)}\| \leq T_h^{-1/2} \|\beta'_{[I_k]} U_{[I_k]} F' b_{j2}\| + \|\beta'_{[I_k]} U_{[I_k]}\| \|T_h^{-1/2} F' b_{j2} - \hat{\xi}_{(j)}\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}.$$

(ii) Assumptions 1, 2 and 3 imply that $\|X_{[I_k]}\| \leq \|\beta_{[I_k]}\| \|F\| + \|U_{[I_k]}\| \lesssim_P q^{1/2} N^{1/2} T^{1/2}$. Together with (i) and Lemma 5(iii), we have

$$\left\| \beta'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(j)} \right\| \leq \left\| \beta'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)} \right\| + \sum_{i=1}^{k-1} \left\| \beta_{[I_k]} \right\| \left\| \frac{X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \left\| \zeta'_{(i)} \tilde{U}_{(i)} \hat{\xi}_{(j)} \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}.$$

(iii) With Lemma 6 and Assumption 4, we have

$$\left\| (u_T)'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)} \right\| \leq T_h^{-1/2} \left\| (u_T)'_{[I_k]} U_{[I_k]} F' b_{j2} \right\| + \left\| (u_T)'_{[I_k]} U_{[I_k]} \right\| \left\| T_h^{-1/2} F' b_{j2} - \hat{\xi}_{(j)} \right\| \lesssim_P T^{-1/2} q N + T^{1/2}.$$

(iv) Similar to (ii), with Lemma 5, $\|X_{[I_k]}\| \lesssim_P q^{1/2} N^{1/2} T^{1/2}$, and (iii), we have

$$\left\| (u_T)'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(j)} \right\| \leq \left\| (u_T)'_{[I_k]} U_{[I_k]} \hat{\xi}_{(j)} \right\| + \sum_{i=1}^{k-1} \left\| (u_T)'_{[I_k]} \right\| \left\| \frac{X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \left\| \zeta'_{(i)} \tilde{U}_{(i)} \hat{\xi}_{(j)} \right\| \lesssim_P T^{-1/2} q N + T^{1/2}.$$

□

Lemma 8. Under assumptions of Theorem 1, for $k \leq \tilde{K}$, we have

- (i) $\|\hat{\varsigma}_{(k)} - \varsigma_{(k)}\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$.
- (ii) $q^{-1/2} N^{-1/2} \left\| \zeta'_{(k)} \beta_{[I_k]} - \varsigma'_{(k)} \beta_{[I_k]} \right\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$.
- (iii) $|\zeta'_{(k)} (u_T)_{[I_k]} - \varsigma'_{(k)} (u_T)_{[I_k]}| \lesssim_P T^{-1} q^{1/2} N^{1/2} + q^{-1/2} N^{-1/2}$.
- (iv) $\left\| \zeta'_{(k)} \tilde{D}_{(k)} - \varsigma'_{(k)} D_{(k)} \right\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$.
- (v) $q^{-1/2} N^{-1/2} \left\| \zeta'_{(k)} \tilde{D}_{(k)} \beta - \varsigma'_{(k)} D_{(k)} \beta \right\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$.
- (vi) $|\zeta'_{(k)} \tilde{D}_{(k)} u_T - \varsigma'_{(k)} D_{(k)} u_T| \lesssim_P T^{-1} q^{1/2} N^{1/2} + q^{-1/2} N^{-1/2}$.

Proof. We prove (i) - (vi) by induction. Consider the $k = 1$ case. The definitions of $\hat{\varsigma}_{(k)}$ in (12) and $\varsigma_{(k)}$ in Section 3.1 lead to

$$\hat{\varsigma}_{(k)} - \varsigma_{(k)} = T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} (\tilde{D}_{(k)} \beta F \hat{\xi}_{(k)} + \tilde{D}_{(k)} U \hat{\xi}_{(k)}) - \lambda_{(k)}^{-1/2} D_{(k)} \beta b_{(k)}, \quad (\text{C.49})$$

when $k = 1$, as $\tilde{D}_{(1)} = D_{(1)} = (\mathbb{I}_N)_{[I_1]}$, (C.49) becomes

$$\hat{\varsigma}_{(1)} - \varsigma_{(1)} = (T_h^{-1/2} \hat{\lambda}_{(1)}^{-1/2} \beta_{(1)} F \hat{\xi}_{(1)} - \lambda_{(1)}^{-1/2} \beta_{(1)} b_{(1)}) + T_h^{-1/2} \hat{\lambda}_{(1)}^{-1/2} U_{[I_1]} \hat{\xi}_{(1)}. \quad (\text{C.50})$$

As Lemma 3(iii) and (iv) imply that $\left\| T_h^{-1/2} \hat{\lambda}_{(1)}^{-1/2} F \hat{\xi}_{(1)} - \lambda_{(1)}^{-1/2} b_{(1)} \right\| \lesssim_P q^{-1} N^{-1} + T^{-1/2} q^{-1/2} N^{-1/2}$ and $\|\beta_{(1)}\| \lesssim q^{1/2} N^{1/2}$, to prove (i) it is sufficient to show that $T_h^{-1/2} q^{-1/2} N^{-1/2} \|U_{[I_1]}\| \lesssim_P T^{-1/2} + q^{-1/2} N^{-1/2}$, which is given by Assumption 3.

(ii) As $\|\beta_{(1)}\| \lesssim q^{1/2}N^{1/2}$, (ii) is implied by (i).

(iii) Left-multiplying (C.50) by $(u_T)'_{[I_1]}$, as $\left\| (u_T)'_{[I_1]} \beta_{[I_1]} \right\| \lesssim_P q^{1/2}N^{1/2}$, to prove (iii) when $k = 1$, it is sufficient to show that $T_h^{-1/2}q^{-1/2}N^{-1/2} \left\| (u_T)'_{[I_1]} U_{[I_1]} \hat{\xi}_{(1)} \right\| \lesssim_P T^{-1}q^{1/2}N^{1/2} + q^{-1/2}N^{-1/2}$, which is implied by Lemma 7.

(iv)-(vi) are equivalent to (i)-(iii) when $k = 1$ as $\tilde{D}_{(1)} = D_{(1)} = (\mathbb{I}_N)_{[I_1]}$.

Then, we assume that (i) - (vi) hold for $i < k$ and prove (i) - (vi) also hold for k .

(i) Similar to the $k = 1$ case, using (C.49) and Lemma 3(iii)(iv), it is sufficient to show that $T_h^{-1/2}q^{-1/2}N^{-1/2} \left\| \tilde{U}_{(k)} \right\| \lesssim_P T^{-1/2} + q^{-1/2}N^{-1/2}$ and $q^{-1/2}N^{-1/2} \left\| (\tilde{D}_{(k)} - D_{(k)})\beta \right\| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2}$. The first inequality is the same as the $k = 1$ case, which is implied by Lemma 5. As to the second inequality, write

$$(\tilde{D}_{(k)} - D_{(k)})\beta = \sum_{i < k} \left(\frac{\beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} \varsigma'_{(i)} D_{(i)}\beta - \frac{X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \varsigma'_{(i)} \tilde{D}_{(i)}\beta \right).$$

As (v) holds for $i < k$, it is sufficient to show that

$$\left\| \frac{\beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} - \frac{X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \lesssim_P T^{-1/2} + q^{-1/2}N^{-1/2}. \quad (\text{C.51})$$

Plugging $X_{[I_k]} = \beta_{[I_k]}F + U_{[I_k]}$ into (C.51) and using Lemma 3(iii)(iv) again, we only need to show that

$$T_h^{-1/2}q^{-1/2}N^{-1/2} \left\| U_{[I_k]} \hat{\xi}_{(i)} \right\| \lesssim_P q^{-1/2}N^{-1/2} + T^{-1/2},$$

which holds by Assumption 3 and $\left\| \hat{\xi}_{(i)} \right\| = 1$.

(ii) It is implied by (i) as $\|\beta_{[I_k]}\| \lesssim q^{1/2}N^{1/2}$.

(iii) By (C.49), we have

$$\begin{aligned} (u_T)'_{[I_k]} \hat{\varsigma}_{(k)} - (u_T)'_{[I_k]} \varsigma_{(k)} &= (u_T)'_{[I_k]} (T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \tilde{D}_{(k)} \beta F \hat{\xi}_{(k)} - \lambda_{(k)}^{-1/2} D_{(k)} \beta_{(k)} b_{(k)}) \\ &\quad + T_h^{-1/2} \hat{\lambda}_{(k)}^{-1/2} (u_T)'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(k)}. \end{aligned}$$

As in the $k = 1$ case, for the second term, we have $T_h^{-1/2}q^{-1/2}N^{-1/2} \left\| (u_T)'_{[I_k]} \tilde{U}_{(k)} \hat{\xi}_{(k)} \right\| \lesssim_P T^{-1}q^{1/2}N^{1/2} + q^{-1/2}N^{-1/2}$, as is given by Lemma 7(iv). For the first term, similar to the proof of (i), using Lemma 3(iii)(iv), it is sufficient to show that $q^{-1/2}N^{-1/2} \left\| (u_T)'_{[I_k]} (\tilde{D}_{(k)} - D_{(k)})\beta \right\| \lesssim_P T^{-1}q^{1/2}N^{1/2} + q^{-1/2}N^{-1/2}$. Write

$$(u_T)'_{[I_k]} (\tilde{D}_{(k)} - D_{(k)})\beta = \sum_{i < k} \left(\frac{(u_T)'_{[I_k]} \beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} \varsigma'_{(i)} D_{(i)}\beta - \frac{(u_T)'_{[I_k]} (\beta_{[I_k]}F + U_{[I_k]}) \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \varsigma'_{(i)} \tilde{D}_{(i)}\beta \right).$$

As (v) holds for $i < k$, we only need to show that

$$\left\| \frac{(u_T)'_{[I_k]} \beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} - \frac{(u_T)'_{[I_k]} (\beta_{[I_k]} F + U_{[I_k]} \hat{\xi}_{(i)})}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \lesssim_P T^{-1} q^{1/2} N^{1/2} + q^{-1/2} N^{-1/2}. \quad (\text{C.52})$$

Using Lemma 3(iii)(iv) again, it is sufficient to show

$$T_h^{-1/2} q^{-1/2} N^{-1/2} \left\| (u_T)'_{[I_k]} U_{[I_k]} \hat{\xi}_{(i)} \right\| \lesssim_P T^{-1} q^{1/2} N^{1/2} + q^{-1/2} N^{-1/2},$$

which holds by Lemma 7(iii).

(iv) By simple algebra, we have

$$\zeta'_{(k)} \tilde{D}_{(k)} - \zeta'_{(k)} D_{(k)} = (\zeta'_{(k)} - \zeta'_{(k)}) (\mathbb{I}_N)_{[I_k]} + \sum_{i < k} \left(\frac{\zeta'_{(k)} \beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} \zeta'_{(i)} D_{(i)} - \frac{\zeta'_{(k)} X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \zeta'_{(i)} \tilde{D}_{(i)} \right).$$

Using the fact that (i) holds for k , (iii) holds for $i < k$ and (C.51), the proof is completed.

(v) Similarly, we have

$$\zeta'_{(k)} \tilde{D}_{(k)} \beta - \zeta'_{(k)} D_{(k)} \beta = (\zeta'_{(k)} \beta_{[I_k]} - \zeta'_{(k)} \beta_{[I_k]}) + \sum_{i < k} \left(\frac{\zeta'_{(k)} \beta_{[I_k]} b_{(i)}}{\sqrt{\lambda_{(i)}}} \zeta'_{(i)} D_{(i)} \beta - \frac{\zeta'_{(k)} X_{[I_k]} \hat{\xi}_{(i)}}{\sqrt{T_h \hat{\lambda}_{(i)}}} \zeta'_{(i)} \tilde{D}_{(i)} \beta \right)$$

Using the fact that (i) holds for k , (v) holds for $i < k$ and (C.51), the proof is completed.

(vi) Replace $q^{-1/2} N^{-1/2} \beta$ by u_T in the proof of (v), we obtain (vi). □

Lemma 9. *Under assumptions of Theorem 1, for $k \leq \tilde{K} + 1$, we have*

$$(i) \quad \left\| \tilde{Z}_{(k)} F' \right\| \lesssim_P T^{1/2} + T q^{-1} N^{-1}.$$

$$(ii) \quad \left\| \tilde{Z}_{(k)} U'_{[I_0]} \right\| \lesssim_P N_0^{1/2} T^{1/2} + T q^{-1/2} N^{-1/2}.$$

Proof. (i) From the definition (18) of $\tilde{Z}_{(k)}$, we have

$$\tilde{Z}_{(k)} F' = Z F' - \sum_{i=1}^{k-1} Y \hat{\xi}_{(i)} \frac{\zeta'_{(i)} \tilde{U}_{(i)} F'}{\sqrt{T_h \hat{\lambda}_{(i)}}}.$$

Then along with Lemma 5(ii), we have

$$\left\| \tilde{Z}_{(k)} F' \right\| \leq \|Z F'\| + \sum_{i=1}^{k-1} \left\| Y \hat{\xi}_{(i)} \right\| \left\| \frac{\zeta'_{(i)} \tilde{U}_{(i)} F'}{\sqrt{T_h \hat{\lambda}_{(i)}}} \right\| \lesssim_P T^{1/2} + T q^{-1} N^{-1}.$$

(ii) With (18) again, we have

$$\tilde{Z}_{(k)}U'_{[I_0]} = ZU'_{[I_0]} - \sum_{i=1}^{k-1} Y\hat{\xi}_{(i)} \frac{\zeta'_{(i)}\tilde{U}_{(i)}U'_{[I_0]}}{\sqrt{T_h\hat{\lambda}_{(i)}}},$$

which, along with Lemma 5(i) and the assumptions on q , lead to

$$\begin{aligned} \left\| \tilde{Z}_{(k)}U'_{[I_0]} \right\| &\leq \left\| ZU'_{[I_0]} \right\| + \sum_{i=1}^{k-1} \left\| Y\hat{\xi}_{(i)} \right\| \left\| \frac{\zeta'_{(i)}\tilde{U}_{(i)}}{\sqrt{T_h\hat{\lambda}_{(i)}}} \right\| \left\| U_{[I_0]} \right\| \\ &\lesssim_P N_0^{1/2}T^{1/2} + \left(q^{-1/2}N^{-1/2} + T^{-1} \right) \left(N_0^{1/2}T^{1/2} + T \right) \\ &\lesssim_P N_0^{1/2}T^{1/2} + Tq^{-1/2}N^{-1/2}. \end{aligned}$$

□

Lemma 10. *Under assumptions of Theorem 1, B_1 , B_2 defined by (31) satisfy*

$$(i) \quad \|B_1\| \lesssim_P 1, \quad \|B_2\| \lesssim_P 1.$$

$$(ii) \quad \|B'_1B_2 - \mathbb{I}_{\tilde{K}}\| \lesssim_P T^{-1} + q^{-1}N^{-1}.$$

$$(iii) \quad \|B_1 - B_2\| \lesssim_P T^{-1/2} + q^{-1}N^{-1}, \quad \|B_1 - B\| \lesssim_P T^{-1/2} + q^{-1/2}N^{-1/2}.$$

$$(iv) \quad \|B_2B'_2 - \mathbb{I}_K\| \lesssim_P T^{-1/2} + q^{-1}N^{-1} \text{ when } \tilde{K} = K.$$

$$(v) \quad \|B_1B'_2 - \mathbb{I}_K\| \lesssim_P T^{-1} + q^{-1}N^{-1}, \text{ when } \tilde{K} = K.$$

Proof. (i) Using the definition (31) of B_1 and Assumption 1, we have

$$\|b_{k1}\| = \left\| \frac{F\hat{\xi}_{(k)}}{\sqrt{T_h}} \right\| \lesssim_P 1,$$

which leads to $\|B_1\| \lesssim_P 1$.

Using the definition (31) of B_2 , we have

$$\|b_{k2}\| = \left\| \frac{\tilde{\beta}'_{(k)}\hat{\xi}_{(k)}}{\sqrt{\hat{\lambda}_{(k)}}} \right\| \leq q^{-1/2}N^{-1/2} \left\| \tilde{\beta}_{(k)} \right\|. \quad (\text{C.53})$$

Note that

$$\left\| \tilde{\beta}_{(k)} \right\| \leq \left\| \beta_{[I_k]} \right\| + \sum_{i=1}^{k-1} \left\| \frac{X_{[I_k]}\hat{\xi}_{(i)}}{\sqrt{T_h\hat{\lambda}_{(i)}}} \right\| \left\| \zeta'_{(i)}\tilde{\beta}_{(i)} \right\| \lesssim_P q^{1/2}N^{1/2} + \sum_{i=1}^{k-1} \left\| \tilde{\beta}_{(i)} \right\| \quad (\text{C.54})$$

and $\|\tilde{\beta}_{(1)}\| = \|\tilde{\beta}_{[I_1]}\| \lesssim q^{1/2}N^{1/2}$, we have $\|\tilde{\beta}_{(k)}\| \lesssim q^{1/2}N^{1/2}$ by induction. Together with (C.53), we have $\|b_{k2}\| \lesssim_P 1$ and thus $\|B_2\| \lesssim_P 1$.

(ii) By (12) and Lemma 1, we have

$$\delta_{lk} = \hat{\xi}_{(l)}' \hat{\xi}_{(k)} = \frac{\hat{\xi}_{(l)}' F' \tilde{\beta}_{(k)}' \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} + \frac{\hat{\xi}_{(l)}' \tilde{U}_{(k)}' \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} = b_{l1}' b_{k2} + \frac{\hat{\xi}_{(l)}' \tilde{U}_{(k)}' \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}}.$$

By Lemma 5(iii), we have

$$|b_{l1}' b_{k2} - \delta_{lk}| \lesssim_P q^{-1} N^{-1} + T^{-1},$$

and thus $\|B_1' B_2 - \mathbb{I}_{\tilde{K}}\| \lesssim_P q^{-1} N^{-1} + T^{-1}$.

(iii) Using (12) and $\tilde{X}_{(k)} = \tilde{\beta}_{(k)} F + \tilde{U}_{(k)}$, we have

$$F \hat{\xi}_{(k)} = \frac{F F' \tilde{\beta}_{(k)}' \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}} + \frac{F \tilde{U}_{(k)}' \hat{\varsigma}_{(k)}}{\sqrt{T_h \hat{\lambda}_{(k)}}}.$$

By the definitions of b_{k1} and b_{k2} , it becomes

$$b_{k1} = \frac{F F'}{T_h} b_{k2} + \frac{F \tilde{U}_{(k)}' \hat{\varsigma}_{(k)}}{T_h \sqrt{\hat{\lambda}_{(k)}}}. \quad (\text{C.55})$$

With $\|B_2\| \lesssim_P 1$, Assumption 1 and Lemma 5(ii), (C.55) leads to

$$b_{k1} - b_{k2} \lesssim_P T^{-1/2} + q^{-1} N^{-1}.$$

This completes the proof. The second inequality of (iii) comes from Lemma 3(iv) directly.

(iv) When $\tilde{K} = K$, $B_2' B_2$ is a $K \times K$ matrix. By (i), (ii) and (iii), we have

$$\|B_2' B_2 - \mathbb{I}_K\| \leq \|B_1' B_2 - \mathbb{I}_K\| + \|B_1 - B_2\| \|B_2\| \lesssim_P T^{-1/2} + q^{-1} N^{-1}.$$

Since B_2 is a $K \times K$ matrix, we have

$$\|B_2 B_2' - \mathbb{I}_K\| = \max_{1 \leq i \leq K} |\lambda_i(B_2 B_2') - 1| = \max_{1 \leq i \leq K} |\lambda_i(B_2' B_2) - 1| = \|B_2' B_2 - \mathbb{I}_K\| \lesssim_P T^{-1/2} + q^{-1} N^{-1}.$$

(v) With respect to $B_1 B_2'$, we have

$$\sigma_K(B_2) \|B_2 B_1' - \mathbb{I}_K\| \leq \|(B_2 B_1' - \mathbb{I}_K) B_2\| = \|B_2 (B_1' B_2 - \mathbb{I}_K)\| \leq \sigma_1(B_2) \|B_1' B_2 - \mathbb{I}_K\|. \quad (\text{C.56})$$

Since (iv) implies that $\sigma_1(B_2)/\sigma_K(B_2) \lesssim_P 1$ when $\tilde{K} = K$, (ii) and (C.56) yield

$$\|B_1 B'_2 - \mathbb{I}_K\| = \|B_2 B'_1 - \mathbb{I}_K\| \leq \frac{\sigma_1(B_2)}{\sigma_K(B_2)} \|B'_1 B_2 - \mathbb{I}_K\| \lesssim_P T^{-1} + q^{-1} N^{-1}. \quad (\text{C.57})$$

□

Lemma 11. *Under Assumptions 1-5, we have*

- (i) $\left\| T_h^{1/2} \hat{\xi}_{(k)} - b'_{k2} F \right\|_{\text{MAX}} \lesssim_P q^{-1/2} N^{-1/2} (\log T)^{1/2} + T^{-1/2} + q^{-1} N^{-1} T^{1/2}.$
- (ii) $\left\| \hat{\lambda}_{(k)}^{1/2} \hat{\beta}_{(k)} - \beta b_{k1} \right\|_{\text{MAX}} \lesssim_P T^{-1/2} (\log N)^{1/2}.$
- (iii) $\left\| \hat{U} - U \right\|_{\text{MAX}} \lesssim_P q^{-1/2} N^{-1/2} (\log T)^{1/2} + T^{-1/2} (\log NT)^{1/2} + q^{-1} N^{-1} T^{1/2}.$
- (iv) $\max_{i \leq N} T_h^{-1/2} \left\| \hat{U}_{[i]} - U_{[i]} \right\| \lesssim_P T^{-1/2} (\log N)^{1/2} + q^{-1/2} N^{-1/2}.$

Proof. (i) Recall that by (C.48), (12), and (14), we have $T_h^{1/2} \hat{\xi}_{(k)} - b'_{k2} F = \hat{\lambda}_{(k)}^{-1/2} \tilde{\zeta}'_{(k)} \tilde{U}_{(k)}$, and $\hat{\xi}_{(k)} = T^{-1/2} \hat{\lambda}_{(k)}^{-1/2} \tilde{X}_{(k)} \hat{\xi}_{(k)} = T^{-1/2} \hat{\lambda}_{(k)}^{-1/2} X_{[I_k]} \hat{\xi}_{(k)}$. Therefore, we have

$$\left\| T_h^{1/2} \hat{\xi}_{(k)} - b'_{k2} F \right\|_{\text{MAX}} \lesssim_P q^{-1} N^{-1} T^{-1/2} \left(\left\| \hat{\xi}_{(k)} F' \beta'_{[I_k]} \tilde{U}_{(k)} \right\|_{\text{MAX}} + \left\| \hat{\xi}_{(k)} U'_{[I_k]} \tilde{U}_{(k)} \right\|_{\text{MAX}} \right). \quad (\text{C.58})$$

When $k = 1$, $\tilde{U}_{(1)} = U_{[I_1]}$, with $\left\| \beta'_{[I_1]} \tilde{U}_{(1)} \right\|_{\text{MAX}} \lesssim_P q^{1/2} N^{1/2} (\log T)^{1/2}$ from Assumption 4 and $\left\| U_{[I_1]} \right\| \lesssim_P q^{1/2} N^{1/2} + T^{1/2}$ from Assumption 3, we have

$$\begin{aligned} \left\| T_h^{1/2} \hat{\xi}_{(1)} - b'_{12} F \right\|_{\text{MAX}} &\lesssim_P q^{-1} N^{-1} \left\| \beta'_{[I_1]} U_{[I_1]} \right\|_{\text{MAX}} + q^{-1} N^{-1} T^{-1/2} \left\| U'_{[I_1]} U_{[I_1]} \right\| \\ &\lesssim_P q^{-1/2} N^{-1/2} (\log T)^{1/2} + T^{-1/2} + q^{-1} N^{-1} T^{1/2}. \end{aligned}$$

Now suppose that this property holds for $i < k$, then for the first term in (C.58), we have

$$\left\| \beta'_{[I_k]} \tilde{U}_{(k)} \right\|_{\text{MAX}} \lesssim \left\| \beta'_{[I_k]} U_{[I_k]} \right\|_{\text{MAX}} + \sum_{i < k} \left\| \beta_{[I_k]} \right\| \left\| T_h^{-1/2} \hat{\lambda}_{(i)}^{-1/2} X_{[I_k]} \hat{\xi}_{(i)} \right\| \left\| \tilde{\zeta}'_{(i)} \tilde{U}_{(i)} \right\|_{\text{MAX}}.$$

The assumption that (i) holds for $i < k$ implies that

$$\left\| \tilde{\zeta}'_{(i)} \tilde{U}_{(i)} \right\|_{\text{MAX}} = \hat{\lambda}_{(i)}^{1/2} \left\| T_h^{1/2} \hat{\xi}_{(i)} - b'_{i2} F \right\|_{\text{MAX}} \lesssim_P (\log T)^{1/2} + q^{1/2} N^{1/2} T^{-1/2} + q^{-1/2} N^{-1/2}.$$

With $\left\| \beta_{[I_k]} \right\| \lesssim q^{1/2} N^{1/2}$ and $\left\| X_{[I_k]} \right\| \leq \left\| \beta_{[I_k]} \right\| \left\| F \right\| + \left\| U_{[I_k]} \right\| \lesssim_P q^{1/2} N^{1/2} T^{1/2}$ and Assumption 4(ii), we have the first term in (C.58) satisfies

$$\begin{aligned} q^{-1} N^{-1} T^{-1/2} \left\| \hat{\xi}_{(k)} F' \beta'_{[I_k]} \tilde{U}_{(k)} \right\|_{\text{MAX}} &\lesssim q^{-1} N^{-1} T^{-1/2} \left\| \hat{\xi}_{(k)} F' \right\| \left\| \beta'_{[I_k]} \tilde{U}_{(k)} \right\|_{\text{MAX}} \\ &\lesssim_P q^{-1} N^{-1} \left\| \beta'_{[I_k]} U_{[I_k]} \right\|_{\text{MAX}} + \sum_{i < k} q^{-1/2} N^{-1/2} \left\| \tilde{\zeta}'_{(i)} \tilde{U}_{(i)} \right\|_{\text{MAX}} \end{aligned}$$

$$\lesssim_P q^{-1/2} N^{-1/2} (\log T)^{1/2} + T^{-1/2} + q^{-1} N^{-1} T^{1/2}.$$

For the second term in (C.58), we have

$$\begin{aligned} \left\| \widehat{\xi}_{(k)} U'_{[I_k]} \widetilde{U}_{(k)} \right\|_{\text{MAX}} &\leq \left\| \widehat{\xi}_{(k)} U'_{[I_k]} \widetilde{U}_{(k)} \right\| \leq \|U_{[I_k]}\| \left\| \widetilde{U}_{(k)} \right\| \\ &\lesssim q^{1/2} N^{1/2} T^{1/2} (\log T)^{1/2} + T, \end{aligned}$$

where we use Assumption 3 and Lemma 5(i) in the last step. Consequently, (i) also holds for k and this concludes the proof by induction.

(ii) By simple algebra, we have

$$\widehat{\lambda}_{(k)}^{1/2} \widehat{\beta}_{(k)} = T_h^{-1/2} X \widehat{\xi}_{(k)} = \beta b_{k1} + T_h^{-1/2} U \widehat{\xi}_{(k)},$$

which leads to

$$\begin{aligned} \left\| \widehat{\lambda}_{(k)}^{1/2} \widehat{\beta}_{(k)} - \beta b_{k1} \right\|_{\text{MAX}} &\lesssim T_h^{-1} \|U F' b_{k2}\|_{\text{MAX}} + T^{-1} \left\| U(T_h^{1/2} \widehat{\xi}_{(k)} - F' b_{k2}) \right\|_{\text{MAX}} \\ &\lesssim_P T^{-1} \|U F'\|_{\text{MAX}} + T^{-1/2} \|U\|_{\text{MAX}} \left\| T_h^{1/2} \widehat{\xi}_{(k)} - F' b_{k2} \right\| \lesssim_P T^{-1/2} (\log N)^{1/2}, \end{aligned}$$

where we use Assumptions 3, 4, and Lemma 6.

(iii) By triangle inequality, we have

$$\begin{aligned} \left\| \widehat{U} - U \right\|_{\text{MAX}} &= \left\| \sum_{k \leq K} \widehat{\beta}_{(k)} \widehat{F}_{(k)} - \beta F \right\|_{\text{MAX}} \\ &\leq \left\| \beta \left(\sum_{k \leq K} b_{k1} b'_{k2} - \mathbb{I}_K \right) F \right\|_{\text{MAX}} + \sum_{k \leq K} \left\| \widehat{\beta}_{(k)} \widehat{F}_{(k)} - \beta b_{k1} b'_{k2} F \right\|_{\text{MAX}}. \end{aligned}$$

By Assumptions 1, 2 and Lemma 10, the first term satisfies

$$\left\| \beta \left(\sum_{k \leq K} b_{k1} b'_{k2} - \mathbb{I}_K \right) F \right\|_{\text{MAX}} \lesssim \|\beta\|_{\text{MAX}} \|F\|_{\text{MAX}} \|B_1 B'_2 - \mathbb{I}_K\| \lesssim_p (\log T)^{1/2} (T^{-1} + q^{-1} N^{-1}).$$

For the second term, note that by triangle inequality we have

$$\begin{aligned} \left\| \widehat{\beta}_{(k)} \widehat{F}_{(k)} - \beta b_{k1} b'_{k2} F \right\|_{\text{MAX}} &\leq \left\| \widehat{\lambda}_{(k)}^{1/2} \widehat{\beta} - \beta b_{k1} \right\|_{\text{MAX}} \|b'_{k2} F\|_{\text{MAX}} + \|\beta b_{k1}\|_{\text{MAX}} \left\| b'_{k2} F - \widehat{\lambda}_{(k)}^{-1/2} \widehat{F}_{(k)} \right\|_{\text{MAX}} \\ &\quad + \left\| \widehat{\lambda}_{(k)}^{1/2} \widehat{\beta} - \beta b_{k1} \right\|_{\text{MAX}} \left\| b'_{k2} F - \widehat{\lambda}_{(k)}^{-1/2} \widehat{F}_{(k)} \right\|_{\text{MAX}}, \end{aligned}$$

which, together with (i)(ii), conclude the proof.

(iv) Because we have

$$T_h^{-1/2} \left\| \widehat{U}_{[i]} - U_{[i]} \right\| = T_h^{-1/2} \left\| \widehat{\beta}_{[i]} \widehat{F} - \beta_{[i]} F \right\|,$$

it then follows from triangle inequality that

$$\left\| \widehat{\beta}_{[i]} \widehat{F} - \beta_{[i]} F \right\| \leq \left\| \beta_{[i]} (B_1 B'_2 - \mathbb{I}_K) F \right\| + \left\| \beta_{[i]} B_1 \right\| \left\| \widehat{\Lambda}^{-1/2} \widehat{F} - B'_2 F \right\| + \left\| \beta_{[i]} B_1 - \widehat{\beta}_{[i]} \widehat{\Lambda}^{1/2} \right\| \left\| \widehat{\Lambda}^{-1/2} \widehat{F} \right\|.$$

We analyze the three terms on the right-hand side one by one. With $T^{-1/2} \left\| \widehat{\Lambda}^{-1/2} \widehat{F} - B'_2 F \right\| \lesssim_P T^{-1} + q^{-1/2} N^{-1/2}$ from Lemma 6, $\|\beta\|_{\text{MAX}} \lesssim 1$, Lemma 10 and (ii), we have

$$\begin{aligned} \max_i T_h^{-1/2} \left\| \beta_{[i]} (B_1 B'_2 - \mathbb{I}_K) F \right\| &\lesssim T_h^{-1/2} \|\beta\|_{\text{MAX}} \|B_1 B'_2 - \mathbb{I}_K\| \|F\| \lesssim_P q^{-1} N^{-1} + T^{-1}, \\ \max_i T_h^{-1/2} \left\| \beta_{[i]} B_1 \right\| \left\| \widehat{\Lambda}^{-1/2} \widehat{F} - B'_2 F \right\| &\lesssim T_h^{-1/2} \|\beta\|_{\text{MAX}} \left\| \widehat{\Lambda}^{-1/2} \widehat{F} - B'_2 F \right\| \lesssim_P q^{-1/2} N^{-1/2} + T^{-1}, \\ \max_i T_h^{-1/2} \left\| \beta_{[i]} B_1 - \widehat{\beta}_{[i]} \widehat{\Lambda}^{1/2} \right\| \left\| \widehat{\Lambda}^{-1/2} \widehat{F} \right\| &\lesssim T_h^{-1/2} \left\| \beta B_1 - \widehat{\beta} \widehat{\Lambda}^{1/2} \right\|_{\text{MAX}} \|F\| \lesssim_P T^{-1/2} (\log N)^{1/2}. \end{aligned}$$

Consequently, we have the desired bound. $\square$

Lemma 12. *Under Assumptions 1-4, for any $I \subset [N]$, we have the following results:*

- (i) $\left\| T_h^{-1} F \mathbb{M}_{W'} F' - \mathbb{I}_K \right\| \lesssim_P T^{-1/2}, \quad \left\| Z \mathbb{M}_{W'} \right\| \lesssim_P T^{1/2}.$
- (ii) $\left\| F \mathbb{M}_{W'} \right\|_{\text{MAX}} \lesssim_P (\log T)^{1/2}, \quad \left\| Z \mathbb{M}_{W'} \right\|_{\text{MAX}} \lesssim_P (\log T)^{1/2}.$
- (iii) $\left\| \beta'_{[I]} U_{[I]} M_{W'} \right\| \lesssim_P |I|^{1/2} T^{1/2}, \quad \left\| \beta'_{[I]} U_{[I]} M_{W'} \right\|_{\text{MAX}} \lesssim_P |I|^{1/2} (\log T)^{1/2}.$
- (iv) $\left\| \beta'_{[I]} U_{[I]} \mathbb{M}_{W'} F' \right\| \lesssim_P |I|^{1/2} T^{1/2} T^{1/2}, \quad \left\| \beta'_{[I]} U_{[I]} \mathbb{M}_{W'} Z' \right\| \lesssim_P |I|^{1/2} T^{1/2} T^{1/2}.$
- (v) $\left\| U \mathbb{M}_{W'} \right\|_{\text{MAX}} \lesssim_P (\log NT)^{1/2}.$
- (vi) $\left\| U \mathbb{M}_{W'} F' \right\|_{\text{MAX}} \lesssim_P (\log N)^{1/2} T^{1/2}, \quad \left\| U \mathbb{M}_{W'} Z' \right\|_{\text{MAX}} \lesssim_P (\log N)^{1/2} T^{1/2}.$
- (vii) $\left\| U_{[I]} \mathbb{M}_{W'} \right\| \lesssim_P |I|^{1/2} + T^{1/2}, \quad \left\| U_{[I]} \mathbb{M}_{W'} F' \right\| \lesssim_P |I|^{1/2} T^{1/2}, \quad \left\| U_{[I]} \mathbb{M}_{W'} Z' \right\| \lesssim_P |I|^{1/2} T^{1/2}.$
- (viii) $\left\| F \mathbb{M}_{W'} Z' \right\| \lesssim_P T^{1/2}, \quad \left\| F \mathbb{M}_{W'} Z' - F Z' \right\| \lesssim_P 1.$
- (ix) $\left\| (u_T)'_{[I]} \underline{U}_{[I]} \mathbb{M}_{W'} F' \right\| \lesssim_P |I| + |I|^{1/2} T^{1/2}, \quad \left\| (u_T)'_{[I]} \underline{U}_{[I]} \mathbb{M}_{W'} Z' \right\| \lesssim_P |I| + |I|^{1/2} T^{1/2}.$

Proof. (i) With $\left\| (WW')^{-1} \right\| \lesssim_P T^{-1}$ from Assumption 1, $\|WF'\| \lesssim_P T^{1/2}$, we have

$$\begin{aligned} \left\| T_h^{-1} F \mathbb{M}_{W'} F' - \mathbb{I}_K \right\| &\leq \left\| T_h^{-1} F F' - \mathbb{I}_K \right\| + \left\| T_h^{-1} F W' (W W')^{-1} W F' \right\| \\ &\leq \left\| T_h^{-1} F F' - \mathbb{I}_K \right\| + T_h^{-1} \|F W'\|^2 \left\| (W W')^{-1} \right\| \lesssim_P T^{-1/2}. \end{aligned}$$

(ii) Using the bound on $\|F\|_{\text{MAX}}$ and the fact that $\|W\| \lesssim T^{1/2}$ by Assumption 1, we have

$$\|F\mathbb{M}_{W'}\|_{\text{MAX}} \leq \|F\|_{\text{MAX}} + \|FW'(WW')^{-1}W\| \leq \|F\|_{\text{MAX}} + \|FW'\| \|(WW')^{-1}\| \|W\| \lesssim_P (\log T)^{1/2}$$

Replacing F by Z in the above proof, we obtain the second inequality.

(iii) Using Assumption 4(ii) and $\|\mathbb{M}_{W'}\| \leq 1$, the first equation holds directly. In addition, we have

$$\left\| \beta'_{[I]} U_{[I]} \mathbb{M}_{W'} \right\|_{\text{MAX}} \lesssim \left\| \beta'_{[I]} U_{[I]} \right\|_{\text{MAX}} + \left\| \beta'_{[I]} U_{[I]} W' \right\| \|(WW')^{-1}\| \|W\| \lesssim_P |I|^{1/2} (\log T)^{1/2}$$

where we use Assumption 1 and Assumption 4(ii) in the last equality.

(iv) With $\|(WW')^{-1}\| \lesssim_P T^{-1}$ from Assumption 1, $\|WF'\| \lesssim_P T^{1/2}$, and by Assumption 4, we have

$$\begin{aligned} \left\| \beta'_{[I]} U_{[I]} \mathbb{M}_{W'} F' \right\| &\leq \left\| \beta'_{[I]} U_{[I]} F' \right\| + \left\| \beta'_{[I]} U_{[I]} \mathbb{P}_{W'} F' \right\| \\ &\leq \left\| \beta'_{[I]} U_{[I]} F' \right\| + \left\| \beta'_{[I]} U_{[I]} W' \right\| \|(WW')^{-1}\| \|WF'\| \lesssim_P |I|^{1/2} T^{1/2}. \end{aligned}$$

Replacing F by Z in the above proof, we have the second inequality in (iv).

(v) Similar to (ii), using Assumption 3 and 4, we have

$$\|U\mathbb{M}_{W'}\|_{\text{MAX}} \lesssim \|U\|_{\text{MAX}} + \|UW'\|_{\text{MAX}} \|(WW')^{-1}W\| \lesssim_P (\log N)^{1/2} + (\log T)^{1/2}.$$

(vi) Similar to (iv), by Assumptions 1 and 3, we have

$$\|U\mathbb{M}_{W'} F'\|_{\text{MAX}} \lesssim \|UF'\|_{\text{MAX}} + \|UW'\|_{\text{MAX}} \|(WW')^{-1}WF'\| \lesssim_P (\log N)^{1/2} T^{1/2}.$$

Replacing F by Z in the above inequality, we also have

$$\|U\mathbb{M}_{W'} Z'\|_{\text{MAX}} \lesssim_P (\log N)^{1/2} T^{1/2}.$$

(vii) With Assumption 3 and $\|\mathbb{M}_{W'}\| \leq 1$, the first inequality holds directly. By Assumptions 1 and 4, we have

$$\|U_{[I]} \mathbb{M}_{W'} F'\| \leq \|U_{[I]} F'\| + \|U_{[I]} W'\| \|(WW')^{-1}\| \|WF'\| \lesssim_P |I|^{1/2} T^{1/2}.$$

Replacing F by Z in the above proof, we also have the third inequality.

(viii) Using Assumption 1 and $\|\mathbb{M}_{W'}\| \leq 1$, we have $\|Z\mathbb{M}_{W'}\| \lesssim_P T^{1/2}$. Using Assumption 1, we have

$$\|F\mathbb{M}_{W'} Z' - FZ'\| = \|F\mathbb{P}_{W'} Z'\| \leq \|FW'\| \|(WW')^{-1}\| \|WZ'\| \lesssim_P 1.$$

Consequently, $\|F\mathbb{M}_{W'}Z'\| \leq \|F\mathbb{M}_{W'}Z' - FZ'\| + \|FZ'\| \lesssim_P T^{1/2}$ as we have $\|FZ'\| \lesssim_P T^{1/2}$ from Assumption 1.

(ix) By Assumptions 1 and 4, we have

$$\left\| (u_T)'_{[I]} \underline{U}_{[I]} \mathbb{M}_{W'} F' \right\| \leq \left\| (u_T)'_{[I]} \underline{U}_{[I]} F' \right\| + \left\| (u_T)'_{[I]} \underline{U}_{[I]} W' \right\| \|(WW')^{-1}\| \|WF'\| \lesssim_P |I| + |I|^{1/2} T^{1/2}.$$

Replacing F by Z , we have the second inequality. $\square$

Lemma 13. For any $N \times K$ matrix β , if $\|T_h^{-1}FF' - \mathbb{I}_K\| \lesssim_P T^{-1/2}$, we have

(i) $\sigma_j(\beta F)/\sigma_j(\beta) = T_h^{1/2} + O_P(1)$ for $j \leq K$.

(ii) If $\sigma_1(\beta) - \sigma_2(\beta) \asymp \sigma_1(\beta)$, then $\left\| \mathbb{P}_{\tilde{\xi}} - T_h^{-1}F'\mathbb{P}_bF \right\| \lesssim_P T^{-1/2}$, where b and $\tilde{\xi}$ are the first right singular vectors of β and βF , respectively.

Proof. (i) For $j \leq K$, $\sigma_j(\beta F)^2 = \lambda_j(\beta FF' \beta') = \lambda_j(\beta' \beta FF')$, which implies

$$\lambda_j(\beta' \beta) \lambda_p(FF') \leq \sigma_j(\beta F)^2 \leq \lambda_j(\beta' \beta) \lambda_1(FF').$$

With the assumption $\|T_h^{-1}FF - \mathbb{I}_K\| \lesssim_P T^{-1/2}$, we have $T_h^{-1/2} \sigma_j(\beta F)/\sigma_j(\beta) = 1 + O_P(T^{-1/2})$ by Weyl's inequality.

(ii) Let ς and $\tilde{\varsigma}$ be the first left singular vectors of β and βF , respectively. Equivalently, ς and $\tilde{\varsigma}$ are the eigenvectors of $\beta\beta'$ and $T_h^{-1}\beta FF' \beta'$. Since $\|\beta\beta' - T_h^{-1}\beta FF' \beta'\| \leq \|\beta\|^2 \|T_h^{-1}FF' - \mathbb{I}_K\| \lesssim_P \sigma_1(\beta)^2 T^{-1/2}$ and $\sigma_1(\beta) - \sigma_2(\beta) \asymp \sigma_1(\beta)$, by sin-theta theorem we have

$$\|\varsigma\varsigma' - \tilde{\varsigma}\tilde{\varsigma}'\| \lesssim \frac{\|\beta\beta' - T_h^{-1}\beta FF' \beta'\|}{\sigma_1(\beta)^2 - \sigma_2(\beta)^2 - O(\|\beta\beta' - T_h^{-1}\beta FF' \beta'\|)} \lesssim_P T^{-1/2}.$$

Using the relationship between left and right singular vectors, we have

$$b' = \frac{\varsigma' \beta}{\sigma_1(\beta)}, \quad \tilde{\xi}' = \frac{\tilde{\varsigma}' \beta F}{\|\beta F\|}.$$

Therefore,

$$\left\| \mathbb{P}_{\tilde{\xi}} - \frac{\sigma_1(\beta)^2}{\|\beta F\|^2} F' \mathbb{P}_b F \right\| = \left\| \tilde{\xi} \tilde{\xi}' - \frac{F' \beta' \varsigma \varsigma' \beta F}{\|\beta F\|^2} \right\| = \left\| \frac{F' \beta' \tilde{\varsigma} \tilde{\varsigma}' \beta F}{\|\beta F\|^2} - \frac{F' \beta' \varsigma \varsigma' \beta F}{\|\beta F\|^2} \right\| \lesssim_P T^{-1/2}. \quad (\text{C.59})$$

By Weyl's inequality, we have $T_h^{-1} \|\beta F\|^2 = \lambda_1(T^{-1}\beta FF' \beta') = \lambda_1(\beta\beta') + O_P(\sigma_1(\beta)^2 T^{-1/2}) = \sigma_1(\beta)^2 + O_P(\sigma_1(\beta)^2 T^{-1/2})$. Plugging this result into (C.59), we have $\left\| \mathbb{P}_{\tilde{\xi}} - T_h^{-1}F' \mathbb{P}_b F \right\| \lesssim_P T^{-1/2}$. $\square$

References

- Bai, J. and S. Ng (2008). Forecasting economic time series using targeted predictors. *Journal of Econometrics* 146(2), 304–317.
- Davis, C. and W. M. Kahan (1970). The rotation of eigenvectors by a perturbation. iii. *SIAM Journal on Numerical Analysis* 7(1), 1–46.
- Fan, J., Y. Liao, and M. Mincheva (2013). Large covariance estimation by thresholding principal orthogonal complements. *Journal of the Royal Statistical Society, B* 75, 603–680.
- Faust, J. and J. H. Wright (2013). Forecasting inflation. In *Handbook of economic forecasting*, Volume 2, pp. 2–56. Elsevier.
- Genre, V., G. Kenny, A. Meyler, and A. Timmermann (2013). Combining expert forecasts: Can anything beat the simple average? *International Journal of Forecasting* 29(1), 108–121.
- Giglio, S. W. and D. Xiu (2021). Asset pricing with omitted factors. *Journal of Political Economy* 129(7), 1947–1990.
- Huang, D., F. Jiang, K. Li, G. Tong, and G. Zhou (2022). Scaled pca: A new approach to dimension reduction. *Management Science* 68(3), 1678–1695.
- Kelly, B. and S. Pruitt (2015). The three-pass regression filter: A new approach to forecasting using many predictors. *Journal of Econometrics* 186(2), 294–316.
- Marcellino, M., J. H. Stock, and M. W. Watson (2006). A comparison of direct and iterated multistep ar methods for forecasting macroeconomic time series. *Journal of econometrics* 135(1-2), 499–526.
- McCracken, M. W. and S. Ng (2016). Fred-md: A monthly database for macroeconomic research. *Journal of Business & Economic Statistics* 34(4), 574–589.
- Stock, J. H. and M. W. Watson (2002). Macroeconomic forecasting using diffusion indexes. *Journal of Business & Economic Statistics* 20(2), 147–162.
- Wang, W. and J. Fan (2017). Asymptotics of empirical eigenstructure for ultra-high dimensional spiked covariance model. *Annals of Statistics* 45, 1342–1374.
- Wedin, P.-Å. (1972). Perturbation bounds in connection with singular value decomposition. *BIT Numerical Mathematics* 12(1), 99–111.