

PAPER

Analysis and Minimization of l_2 -Sensitivity for Block-State Realization of IIR Digital Filters

Akimitsu DOI^{†a)}, *Member*, Takao HINAMOTO^{††b)}, *Fellow*, and Wu-Sheng LU^{†††c)}, *Nonmember*

SUMMARY Block-state realization of state-space digital filters offers reduced implementation complexity relative to canonical state-space filters while filter's internal structure remains accessible. In this paper, we present a quantitative analysis on l_2 coefficient sensitivity of block-state digital filters. Based on this, we develop two techniques for minimizing average l_2 -sensitivity subject to l_2 -scaling constraints. One of the techniques is based on a Lagrange function and some matrix-theoretic techniques. The other solution method converts the problem at hand into an unconstrained optimization problem which is solved by using an efficient quasi-Newton algorithm where the key gradient evaluation is done in closed-form formulas for fast and accurate execution of quasi-Newton iterations. A case study is presented to demonstrate the validity and effectiveness of the proposed techniques.

key words: block-state realization, l_2 -scaling constraints, l_2 -sensitivity analysis, l_2 -sensitivity minimization, state-space digital filters

1. Introduction

An important issue involved in the implementation of IIR digital filters using fixed-point arithmetic is reducing roundoff noise at the filter's output. As is well known, roundoff noise critically depends on the internal structure of IIR digital filters. In this regard, state-space models for IIR digital filters provide a suitable platform in which the internal structure of a filter can be explored so as to minimize its roundoff noise by choosing an appropriate linear transformation for state-space coordinates without altering the input-output characteristic of the filter. Effective techniques for constructing an optimal filter structure that minimizes the roundoff noise at the filter output subject to l_2 -scaling constraints have been explored [1]–[3]. However, such a realization requires $(n + 1)^2$ multiplications to compute each output sample for an n -th order filter, an increase of n^2 multiplications over canonical direct-form structures. This is a major disadvantage especially for high-order digital filters as its effect on data throughput rate becomes much greater when n is large. Alternatively, block implementation of digital filters was proposed as a method

of increasing data throughput rate through parallelism, and side benefit of using it in terms of decreased roundoff noise due to finite-word-length (FWL) effects was also recognized, see [4]–[6] and the references cited therein. It is true that a block implementation induces additional delay for a given sampling rate. But it is not a concern as long as the system allows a much higher data rate where the sampling period is greatly reduced. So essentially choosing between conventional and block implementations has to do with selecting between sequential (serial) and block (parallel) operations. Parallel processing of block realization schemes makes it well suited for hardware implementations using VLSI circuits as seen in the Texas Instruments TMS320, Motorola 56000, and Analog Devices ADSP-2100 families. Another important and related issue is reducing the effects of quantizing filter coefficients. It is of importance to note that coefficient quantization usually alters filter characteristics. For instance, a stable filter designed under the assumption of infinite precision may become unstable after coefficient quantization. This has motivated the study of coefficient sensitivity and its minimization for digital filters. There are two main classes of techniques for synthesizing an optimal filter structure that minimizes the deviation of a transfer function caused by coefficient quantization, namely l_1/l_2 -sensitivity minimization [7]–[11] and l_2 -sensitivity minimization [12]–[17]. It has been argued [12]–[15] that the sensitivity measure based on the l_2 norm is more natural and reasonable relative to that based on the l_1/l_2 -sensitivity minimization. Alternatively, the problem of minimizing the l_2 -sensitivity subject to l_2 -scaling constraints for state-space digital filters has been studied to suppress overflow oscillations [16], [17]. The analysis and minimization of the l_1/l_2 -sensitivity for the block-state realization of IIR digital filters has also been examined in [18], [19]. A bit earlier, our study on the analysis of l_2 -sensitivity and its minimization subject to l_2 -scaling constraints for block-state realization of IIR digital filters has been included in [20, Sec. 17.4].

This paper presents a quantitative analysis on l_2 -sensitivity for block-state realization of state-space digital filters [20, Sec. 17.4.1]. Following the analysis, we develop two techniques for minimizing a sensitivity measure known as average l_2 -sensitivity subject to l_2 -scaling constraints for the block-state realization of state-space digital filters [20, Sec. 17.4.2]. One of the techniques developed is based on a Lagrange function, while the other relies on an efficient quasi-Newton algorithm. Unlike the technique based on a Lagrange function in [17], [20, Sec. 17.4.2.1], the present

Manuscript received August 8, 2017.

Manuscript revised September 19, 2017.

[†]The author is with Faculty of Applied Information Science, Hiroshima Institute of Technology, Hiroshima-shi, 731-5193 Japan.

^{††}The author is with Hiroshima University, Higashi-hiroshima-shi, 739-8527 Japan.

^{†††}The author is with the Department of Electrical and Computer Engineering, University of Victoria, Victoria, B.C., Canada, V8W 3P6.

a) E-mail: doi@cc.it-hiroshima.ac.jp

b) E-mail: hinamoto@ieee.org

c) E-mail: wslu@ece.uvic.ca

DOI: 10.1587/transfun.E101.A.447

paper explores a much easy-to-implement algorithmic steps and proposes a warm-start initial estimate that helps accelerate the optimization process. In [16], the quasi-Newton algorithm was applied where the l_2 -sensitivity measure and its gradients are partially expressed by infinite sums, thus these terms can be evaluated only approximately. In the present paper, closed-form formulas for the l_2 -sensitivity measure and its gradients are derived for improved accuracy and speed. Experimental results are presented to demonstrate the validity and effectiveness of the proposed methods. This paper can be viewed as a more detailed and partially improved version of the study in [20, Sec. 17.4].

Throughout, $\mathbf{I}_n$ denotes the identity matrix of dimension $n \times n$, $\mathbf{A}^T$ ($\mathbf{A}^H$) denotes transpose (conjugate transpose) of a matrix $\mathbf{A}$, and $\text{tr}[\mathbf{A}]$ denotes the trace of square matrix $\mathbf{A}$. Bold uppercase, bold lowercase and plain lowercase are used to make the distinction between matrices, vectors and scalar values, respectively.

2. l_2 -Sensitivity Analysis and Problem Formulation

A. Block-state realization of a state-space model

Consider a stable, controllable and observable state-space model $(\mathbf{A}, \mathbf{b}, \mathbf{c}, d)_n$ described by

$$\begin{aligned} \mathbf{x}(k+1) &= \mathbf{A}\mathbf{x}(k) + \mathbf{b}u(k) \\ y(k) &= \mathbf{c}\mathbf{x}(k) + du(k) \end{aligned} \quad (1)$$

where $\mathbf{x}(k)$ is an $n \times 1$ state vector, $u(k)$ is a scalar input, $y(k)$ is a scalar output, and $\mathbf{A}, \mathbf{b}, \mathbf{c}$ and d are $n \times n, n \times 1, 1 \times n$, and 1×1 real constant matrices, respectively. Note that in the most pessimistic case where all the elements of coefficient matrices in (1) are nontrivial, $(n+1)^2$ multiplications are required to compute the output $y(k)$ during each sample interval.

Equation (1) can be written as [4]–[6]

$$\begin{aligned} \hat{\mathbf{x}}(k+1) &= \hat{\mathbf{A}}\hat{\mathbf{x}}(k) + \hat{\mathbf{B}}\hat{\mathbf{u}}(k) \\ \hat{\mathbf{y}}(k) &= \hat{\mathbf{C}}\hat{\mathbf{x}}(k) + \hat{\mathbf{D}}\hat{\mathbf{u}}(k) \end{aligned} \quad (2a)$$

where $\hat{\mathbf{x}}(k) = \mathbf{x}(kL)$

$$\begin{aligned} \hat{\mathbf{u}}(k) &= \begin{bmatrix} u(kL) \\ u(kL+1) \\ \vdots \\ u(kL+L-1) \end{bmatrix}, \quad \hat{\mathbf{y}}(k) = \begin{bmatrix} y(kL) \\ y(kL+1) \\ \vdots \\ y(kL+L-1) \end{bmatrix} \\ \hat{\mathbf{A}} &= \mathbf{A}^L, \quad \hat{\mathbf{B}} = [\mathbf{A}^{L-1}\mathbf{b} \quad \mathbf{A}^{L-2}\mathbf{b} \quad \dots \quad \mathbf{b}] \\ \hat{\mathbf{C}} &= \begin{bmatrix} \mathbf{c} \\ \mathbf{c}\mathbf{A} \\ \vdots \\ \mathbf{c}\mathbf{A}^{L-1} \end{bmatrix}, \quad \hat{\mathbf{D}} = \begin{bmatrix} d & 0 & \dots & 0 \\ \mathbf{c}\mathbf{b} & d & \ddots & \vdots \\ \vdots & \ddots & \ddots & 0 \\ \mathbf{c}\mathbf{A}^{L-2}\mathbf{b} & \dots & \mathbf{c}\mathbf{b} & d \end{bmatrix} \end{aligned} \quad (2b)$$

From (2a), it follows that a total of

$$m = n^2 + 2nL + \frac{L(L+1)}{2} \quad (3)$$

multiplications are required in each block of L output samples or, equivalently, an average of

$$\frac{m}{L} = \frac{n^2}{L} + 2n + \frac{L+1}{2} \quad (4)$$

multiplications for each output sample. If the average measure, m/L , in (4) is minimized with respect to L , the optimal block length is obtained as

$$L = \sqrt{2n} \quad (5)$$

which is non-integer and requires rounding to the nearest integer. By substituting (5) into (4), the minimum value of m/L becomes

$$\left(\frac{m}{L}\right)_{\min} = (2 + \sqrt{2})n + \frac{1}{2} \quad (6)$$

This compares favorably with the processing complexity for the canonical forms of the SISO state-space model in (1), which requires $(2n+1)$ multiplications per output sample, and reveals that the system in (2a) enables us to perform fast processing for high-order digital filters.

In the rest of the paper, the system described by (2a) is referred to as *block-state realization*, $(\hat{\mathbf{A}}, \hat{\mathbf{B}}, \hat{\mathbf{C}}, \hat{\mathbf{D}})_n$, that is generated by the state-space model $(\mathbf{A}, \mathbf{b}, \mathbf{c}, d)_n$ in (1). In this way, (2b) defines a mapping $(\mathbf{A}, \mathbf{b}, \mathbf{c}, d)_n \rightarrow (\hat{\mathbf{A}}, \hat{\mathbf{B}}, \hat{\mathbf{C}}, \hat{\mathbf{D}})_n$ that transforms the state-space model in (1) to the block-state realization in (2).

The block-state realization was proposed to achieve a high speed algorithm for the state-space model in (1) [4]–[6]. The system in (2a) can be realized by an L -input/ L -output digital filter implemented with serial-in/parallel-out and parallel-in/serial-out registers, as shown in Fig. 1. Hence, the number of multiplications per output sample can be drastically reduced by updating the state vector that requires great many multiplications once in L samples. Moreover, because a block-state realization preserves the properties of the state-space model [5], the robust block-state realization can be constructed from the state-space model that minimizes FWL effects. In other words, we can realize the filtering with high speed and accuracy (due to the reduced number of computations per output sample) via the block-state realization. It

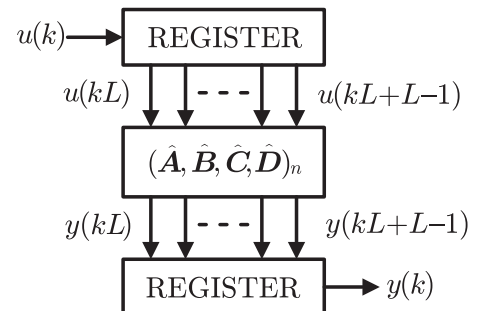


Fig. 1 Block-state realization using serial-in/parallel-out and parallel-in/serial-out registers.

is obvious that the serial input and output sample rates are L times the fundamental clock rate of the system in (2a).

B. Analysis of l_2 -sensitivity

B.1 Definition of sensitivity function

Consider an MIMO transfer function $\mathbf{H}(z)$ that contains N parameters $\{p_1, p_2, \dots, p_N\}$. Let $\{\tilde{p}_i\}$ be the FWL version of $\{p_i\}$ where $\tilde{p}_i = p_i + \Delta p_i$ with Δp_i the parameter perturbation, and $\mathbf{H}_{act}(z)$ be the transfer function associated with perturbed parameters $\{\tilde{p}_i\}$. The first-order approximation of $\mathbf{H}_{act}(z)$ then gives

$$\mathbf{H}_{act}(z) = \mathbf{H}(z) + \Delta \mathbf{H}(z) \quad (7a)$$

where $\Delta \mathbf{H}(z)$ will be

$$\Delta \mathbf{H}(z) = \sum_{i=1}^N \frac{\partial \mathbf{H}(z)}{\partial p_i} \Delta p_i \quad (7b)$$

For a fixed-point implementation of B bits, the parameter perturbations are considered to be independent uniformly distributed random variables within a range $[-2^{-B-1}, 2^{-B-1}]$. Then, a measure of the deviation in the transfer function can statistically be defined as

$$\begin{aligned} \sigma_{\Delta H}^2 &= \frac{1}{2\pi} \int_0^{2\pi} E(\text{tr}[\Delta \mathbf{H}^H(e^{j\omega}) \Delta \mathbf{H}(e^{j\omega})]) d\omega \\ &= \frac{1}{2\pi j} \oint_{|z|=1} E(\text{tr}[\Delta \mathbf{H}^H(z) \Delta \mathbf{H}(z)]) \frac{dz}{z} \end{aligned} \quad (8)$$

where $E(\cdot)$ denotes the ensemble average operation. Since $\{\Delta p_i\}$ are independent uniformly distributed random variables, it follows that

$$E(\text{tr}[\Delta \mathbf{H}^H(z) \Delta \mathbf{H}(z)]) = \sum_{i=1}^N \text{tr} \left[\frac{\partial \mathbf{H}^H(z)}{\partial p_i} \cdot \frac{\partial \mathbf{H}(z)}{\partial p_i} \right] \sigma^2 \quad (9)$$

where $\sigma^2 = E((\Delta p_i)^2) = 2^{-2B}/12$. Equation (9) establishes an analytic relationship between deviation in the transfer function induced by an FWL block-state realization and parameter sensitivity.

We are now in a position to define the l_2 -sensitivity of the block-state realization in (2a).

Definition 1: Let $\mathbf{X}$ be an $m \times n$ real matrix and let $f(\mathbf{X})$ be a scalar complex function of $\mathbf{X}$, differentiable with respect to all entries of $\mathbf{X}$. The sensitivity function of $f(\mathbf{X})$ with respect to $\mathbf{X}$ is then defined as

$$\frac{\partial f(\mathbf{X})}{\partial \mathbf{X}} = \begin{bmatrix} \frac{\partial f(\mathbf{X})}{\partial x_{11}} & \dots & \frac{\partial f(\mathbf{X})}{\partial x_{1n}} \\ \vdots & \ddots & \vdots \\ \frac{\partial f(\mathbf{X})}{\partial x_{m1}} & \dots & \frac{\partial f(\mathbf{X})}{\partial x_{mn}} \end{bmatrix} \quad (10)$$

where x_{il} denotes the (i, l) th entry of matrix $\mathbf{X}$.

B.2 l_2 -sensitivity of system (2)

We now apply the above definition to evaluate the sensitivity of the transfer function of the block-state realization in (2a) which is given by

$$\mathbf{H}(z) = \hat{\mathbf{C}}(z\mathbf{I}_n - \hat{\mathbf{A}})^{-1} \hat{\mathbf{B}} + \hat{\mathbf{D}} \quad (11)$$

whose (i, l) th element is described by

$$H_{il}(z) = \hat{\mathbf{c}}_i(z\mathbf{I}_n - \hat{\mathbf{A}})^{-1} \hat{\mathbf{b}}_l + \hat{d}_{il} \quad (12a)$$

where

$$\begin{aligned} \hat{\mathbf{B}} &= [\hat{\mathbf{b}}_1 \quad \hat{\mathbf{b}}_2 \quad \dots \quad \hat{\mathbf{b}}_L] \\ \hat{\mathbf{C}} &= \begin{bmatrix} \hat{\mathbf{c}}_1 \\ \hat{\mathbf{c}}_2 \\ \vdots \\ \hat{\mathbf{c}}_L \end{bmatrix}, \quad \hat{d}_{il} = \begin{cases} 0 & \text{for } i < l \\ d & \text{for } i = l \\ \mathbf{c} \mathbf{A}^{i-l-1} \mathbf{b} & \text{for } i > l \end{cases} \end{aligned} \quad (12b)$$

By virtue of (12a), Definition 1, and the formula

$$\frac{\partial \mathbf{A}^{-1}}{\partial a} = -\mathbf{A}^{-1} \frac{\partial \mathbf{A}}{\partial a} \mathbf{A}^{-1} \quad (13)$$

with a an entry of any invertible matrix $\mathbf{A}$, the sensitivities of $H_{il}(z)$ with respect to matrices $\hat{\mathbf{A}}$, $\hat{\mathbf{b}}_l$, $\hat{\mathbf{c}}_i^T$ and $\hat{d}_{il}$ are found to be

$$\begin{aligned} \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{A}}} &= [\mathbf{f}_l(z) \mathbf{g}_i(z)]^T, \quad \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{b}}_l} = \mathbf{g}_i^T(z) \\ \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{c}}_i^T} &= \mathbf{f}_l(z), \quad \frac{\partial H_{il}(z)}{\partial \hat{d}_{il}} = u_o(i-l) \end{aligned} \quad (14a)$$

respectively, where

$$\begin{aligned} \mathbf{f}_l(z) &= (z\mathbf{I}_n - \hat{\mathbf{A}})^{-1} \hat{\mathbf{b}}_l = (z\mathbf{I}_n - \mathbf{A}^L)^{-1} \mathbf{A}^{L-l} \mathbf{b} \\ &= \sum_{k=0}^{\infty} \mathbf{A}^{(k+1)L-l} \mathbf{b} z^{-(k+1)} \\ \mathbf{g}_i(z) &= \hat{\mathbf{c}}_i(z\mathbf{I}_n - \hat{\mathbf{A}})^{-1} = \mathbf{c} \mathbf{A}^{i-1} (z\mathbf{I}_n - \mathbf{A}^L)^{-1} \\ &= \sum_{k=0}^{\infty} \mathbf{c} \mathbf{A}^{kL+i-1} z^{-(k+1)} \\ u_o(i) &= \begin{cases} 1 & \text{for } i \geq 0 \\ 0 & \text{for } i < 0 \end{cases} \end{aligned} \quad (14b)$$

In the rest of this section, $\mathbf{f}_l(z)$ and $\mathbf{g}_i(z)$ are referred to as *intermediate functions*.

Definition 2: Let $\mathbf{X}(z)$ be an $m \times n$ complex matrix-valued function of a complex variable z and let $x_{pq}(z)$ be the (p, q) th entry of $\mathbf{X}(z)$. The l_2 -norm of $\mathbf{X}(z)$ is then defined as

$$\begin{aligned} \|\mathbf{X}(z)\|_2 &= \left[\frac{1}{2\pi} \int_0^{2\pi} \sum_{p=1}^m \sum_{q=1}^n |x_{pq}(e^{j\omega})|^2 d\omega \right]^{\frac{1}{2}} \\ &= \left(\text{tr} \left[\frac{1}{2\pi j} \oint_{|z|=1} \mathbf{X}(z) \mathbf{X}^H(z) \frac{dz}{z} \right] \right)^{\frac{1}{2}} \end{aligned} \quad (15)$$

From (14) and Definition 2, the l_2 -sensitivity of the subsystem (12a) can be expressed as

$$\begin{aligned} S_{il} &= \left\| \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{A}}} \right\|_2^2 + \left\| \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{b}}_l} \right\|_2^2 + \left\| \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{c}}_i^T} \right\|_2^2 + \left\| \frac{\partial H_{il}(z)}{\partial \hat{\mathbf{d}}_{il}} \right\|_2^2 \\ &= \left\| [\mathbf{f}_l(z) \mathbf{g}_i(z)]^T \right\|_2^2 + \left\| \mathbf{g}_i^T(z) \right\|_2^2 + \left\| \mathbf{f}_l(z) \right\|_2^2 + u_o(i-l) \end{aligned} \quad (16)$$

In order to evaluate the integrals involved in (16), we utilize *Cauchy's integral theorem*

$$\frac{1}{2\pi j} \oint_C z^i \frac{\partial z}{z} = \begin{cases} 1, & i = 0 \\ 0, & i \neq 0 \end{cases} \quad (17)$$

with C a counterclockwise contour that encircles the origin. Then by virtue of (15), (16) becomes

$$\begin{aligned} S_{il} &= \text{tr}[\mathbf{N}_{il}] + \sum_{k=0}^{\infty} \text{tr}[(\mathbf{c} \mathbf{A}^{kL+i-1})^T \mathbf{c} \mathbf{A}^{kL+i-1}] \\ &\quad + \sum_{k=0}^{\infty} \text{tr}[\mathbf{A}^{(k+1)L-l} \mathbf{b} (\mathbf{A}^{(k+1)L-l} \mathbf{b})^T] + u_o(i-l) \end{aligned} \quad (18a)$$

where $\text{tr}[\mathbf{N}_{il}] = \text{tr}[\mathbf{M}_{il}]$ and

$$\begin{aligned} \mathbf{N}_{il} &= \frac{1}{2\pi j} \oint_{|z|=1} [\mathbf{f}_l(z) \mathbf{g}_i(z)]^T \mathbf{f}_l(z^{-1}) \mathbf{g}_i(z^{-1}) \frac{dz}{z} \\ \mathbf{M}_{il} &= \frac{1}{2\pi j} \oint_{|z|=1} \mathbf{f}_l(z^{-1}) \mathbf{g}_i(z^{-1}) [\mathbf{f}_l(z) \mathbf{g}_i(z)]^T \frac{dz}{z} \end{aligned} \quad (18b)$$

A closed-form solution for evaluating matrices $\mathbf{N}_{il}$ and $\mathbf{M}_{il}$ will be deduced shortly.

We note that each output $y(kL+i-1)$ for $i = 1, 2, \dots, L$ in the output vector $\mathbf{y}(k)$ is generated by the subsystem

$$\begin{aligned} \mathbf{H}_i(z) &= [\mathbf{H}_{i1}(z), \mathbf{H}_{i2}(z), \dots, \mathbf{H}_{iL}(z)] \\ &= \hat{\mathbf{c}}_i (z\mathbf{I} - \hat{\mathbf{A}})^{-1} \hat{\mathbf{B}} + \hat{\mathbf{d}}_i \end{aligned} \quad (19)$$

where $\hat{\mathbf{d}}_i = [\hat{d}_{i1}, \hat{d}_{i2}, \dots, \hat{d}_{iL}]$. From this in conjunction with (18a), the l_2 -sensitivity measure for the subsystem in (19) is found to be

$$\begin{aligned} S_i &= \sum_{l=1}^L S_{il} \\ &= \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}] + \sum_{l=1}^L \sum_{k=0}^{\infty} \text{tr}[(\mathbf{c} \mathbf{A}^{kL+i-1})^T \mathbf{c} \mathbf{A}^{kL+i-1}] \\ &\quad + \sum_{k=0}^{\infty} \text{tr}[\mathbf{A}^k \mathbf{b} (\mathbf{A}^k \mathbf{b})^T] + i \quad \text{for } i = 1, 2, \dots, L \end{aligned} \quad (20)$$

Hence, the overall l_2 -sensitivity for the block-state realization

in (11) can be written as

$$\begin{aligned} S &= \sum_{i=1}^L S_i = \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}] + \sum_{l=1}^L \sum_{k=0}^{\infty} \text{tr}[(\mathbf{c} \mathbf{A}^k)^T \mathbf{c} \mathbf{A}^k] \\ &\quad + \sum_{i=1}^L \sum_{k=0}^{\infty} \text{tr}[\mathbf{A}^k \mathbf{b} (\mathbf{A}^k \mathbf{b})^T] + \frac{L(L+1)}{2} \end{aligned} \quad (21)$$

which is equivalent to

$$\begin{aligned} S &= \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}] + L \text{tr}[\mathbf{W}_o] + L \text{tr}[\mathbf{K}_c] \\ &\quad + \frac{L(L+1)}{2} \end{aligned} \quad (22a)$$

where $\mathbf{K}_c$ and $\mathbf{W}_o$ are the controllability and observability Grammians of the state-space model in (1), defined by

$$\mathbf{K}_c = \sum_{k=0}^{\infty} \mathbf{A}^k \mathbf{b} (\mathbf{A}^k \mathbf{b})^T, \quad \mathbf{W}_o = \sum_{k=0}^{\infty} (\mathbf{c} \mathbf{A}^k)^T \mathbf{c} \mathbf{A}^k \quad (22b)$$

that can be obtained by solving the Lyapunov equations

$$\mathbf{K}_c = \mathbf{A} \mathbf{K}_c \mathbf{A}^T + \mathbf{b} \mathbf{b}^T, \quad \mathbf{W}_o = \mathbf{A}^T \mathbf{W}_o \mathbf{A} + \mathbf{c}^T \mathbf{c} \quad (22c)$$

The l_2 -sensitivity measure for the state-space model in (1) was derived in [17], namely,

$$S_o = \text{tr}[\mathbf{N}_{11}] + \text{tr}[\mathbf{W}_o] + \text{tr}[\mathbf{K}_c] + 1 \quad (23)$$

The l_2 -sensitivity of a conventional state-space model, S_o , can be regarded as a special case of the l_2 -sensitivity of a block realization with $L = 1$. However, for the general case with $L > 1$, the two sensitivity measures S and S_o are not directly comparable because the realization scheme has changed from serial to parallel. The average sensitivity S_{ave} comes as a reasonable measure as comparing it with S_o becomes meaningful and it has a straightforward connection to sensitivity S by a scaling factor of $1/L$. That is

$$S_{ave} = \frac{1}{L} \sum_{i=1}^L S_i = \frac{S}{L} \quad (24)$$

which in conjunction with (22a) gives

$$\begin{aligned} S_{ave} &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}] + \text{tr}[\mathbf{W}_o] + \text{tr}[\mathbf{K}_c] \\ &\quad + \frac{L+1}{2} \end{aligned} \quad (25)$$

Notice that the same idea was applied to evaluate the roundoff noise [4] and the l_1/l_2 -sensitivity [18], [19] in the block-state realization.

C. Problem formulation

To identify an optimal internal structure of a given IIR filter that achieves the minimum average l_2 -sensitivity, consider a coordinate transformation defined by

$$\bar{\mathbf{x}}(k) = \mathbf{T}^{-1} \mathbf{x}(k) \quad (26)$$

that is applied to the state-space model in (1). The new realization associated with state vector $\bar{\mathbf{x}}(k)$, denoted by $(\bar{\mathbf{A}}, \bar{\mathbf{b}}, \bar{\mathbf{c}}, d)_n$, is related to the original realization as

$$\bar{\mathbf{A}} = \mathbf{T}^{-1} \mathbf{A} \mathbf{T}, \quad \bar{\mathbf{b}} = \mathbf{T}^{-1} \mathbf{b}, \quad \bar{\mathbf{c}} = \mathbf{c} \mathbf{T} \quad (27)$$

For this realization, the new controllability Grammian $\bar{\mathbf{K}}_c$ and new observability Grammian $\bar{\mathbf{W}}_o$ are related to the original Grammians as

$$\bar{\mathbf{K}}_c = \mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T}, \quad \bar{\mathbf{W}}_o = \mathbf{T}^T \mathbf{W}_o \mathbf{T} \quad (28)$$

respectively, the new intermediate functions $\bar{f}_j(z)$ and $\bar{g}_i(z)$ are found to be [14], [17]

$$\bar{f}_l(z) = \mathbf{T}^{-1} f_l(z), \quad \bar{g}_i(z) = g_i(z) \mathbf{T} \quad (29)$$

respectively, and the canonical-state to block-state mapping is given by

$$(\bar{\mathbf{A}}, \bar{\mathbf{b}}, \bar{\mathbf{c}}, d)_n \rightarrow (\mathbf{T}^{-1} \hat{\mathbf{A}} \mathbf{T}, \mathbf{T}^{-1} \hat{\mathbf{B}}, \hat{\mathbf{C}} \mathbf{T}, \hat{D})_n \quad (30)$$

Moreover, matrices N_{il} and M_{il} introduced in (18b) are transformed to $\bar{N}_{il}$ and $\bar{M}_{il}$ as

$$\begin{aligned} \bar{N}_{il} &= \frac{1}{2\pi j} \oint_{|z|=1} [\bar{f}_l(z) \bar{g}_i(z)]^T \bar{f}_l(z^{-1}) \bar{g}_i(z^{-1}) \frac{dz}{z} \\ &= \mathbf{T}^T N_{il}(\mathbf{P}) \mathbf{T} \\ \bar{M}_{il} &= \frac{1}{2\pi j} \oint_{|z|=1} \bar{f}_l(z^{-1}) \bar{g}_i(z^{-1}) [\bar{f}_l(z) \bar{g}_i(z)]^T \frac{dz}{z} \\ &= \mathbf{T}^{-1} M_{il}(\mathbf{P}) \mathbf{T}^{-T} \end{aligned} \quad (31a)$$

respectively, where $\mathbf{P} = \mathbf{T} \mathbf{T}^T$ and

$$\begin{aligned} N_{il}(\mathbf{P}) &= \frac{1}{2\pi j} \oint_{|z|=1} [f_l(z) g_i(z)]^T \mathbf{P}^{-1} f_l(z^{-1}) g_i(z^{-1}) \frac{dz}{z} \\ M_{il}(\mathbf{P}) &= \frac{1}{2\pi j} \oint_{|z|=1} f_l(z^{-1}) g_i(z^{-1}) \mathbf{P} [f_l(z) g_i(z)]^T \frac{dz}{z} \end{aligned} \quad (31b)$$

Hence $N_{il} = N_{il}(\mathbf{I}_n)$ and $M_{il} = M_{il}(\mathbf{I}_n)$ follow from (18b). It is noted that

$$\begin{aligned} \bar{f}_l(z) \bar{g}_i(z) &= \mathbf{T}^{-1} f_l(z) g_i(z) \mathbf{T} \\ &= \begin{bmatrix} \mathbf{T}^{-1} & \mathbf{0} \end{bmatrix} \begin{bmatrix} z\mathbf{I}_n - \mathbf{A}^L & -\mathbf{A}^{L-l} \mathbf{b} \mathbf{c} \mathbf{A}^{i-1} \\ \mathbf{0} & z\mathbf{I}_n - \mathbf{A}^L \end{bmatrix}^{-1} \begin{bmatrix} \mathbf{0} \\ \mathbf{T} \end{bmatrix} \end{aligned} \quad (32)$$

If we denote the observability Grammian of a composite

system $\bar{f}_l(z) \bar{g}_i(z)$ in (32) by $\mathbf{Y}_{il}$, matrix $N_{il}(\mathbf{P})$ can be obtained by solving the Lyapunov equation

$$\begin{aligned} \mathbf{Y}_{il} &= \begin{bmatrix} \mathbf{A}^L & \mathbf{A}^{L-l} \mathbf{b} \mathbf{c} \mathbf{A}^{i-1} \\ \mathbf{0} & \mathbf{A}^L \end{bmatrix}^T \mathbf{Y}_{il} \begin{bmatrix} \mathbf{A}^L & \mathbf{A}^{L-l} \mathbf{b} \mathbf{c} \mathbf{A}^{i-1} \\ \mathbf{0} & \mathbf{A}^L \end{bmatrix} \\ &+ \begin{bmatrix} \mathbf{P}^{-1} & \mathbf{0} \\ \mathbf{0} & \mathbf{0} \end{bmatrix} \end{aligned} \quad (33)$$

and then taking the lower-right $n \times n$ block of $\mathbf{Y}_{il}$ as $N_{il}(\mathbf{P})$:

$$N_{il}(\mathbf{P}) = \begin{bmatrix} \mathbf{0} & \mathbf{I}_n \end{bmatrix} \mathbf{Y}_{il} \begin{bmatrix} \mathbf{0} \\ \mathbf{I}_n \end{bmatrix} \quad (34)$$

Therefore, under coordinate transformation $\bar{\mathbf{x}}(k) = \mathbf{T}^{-1} \mathbf{x}(k)$ in (26), the average l_2 -sensitivity measure in (25) becomes

$$\begin{aligned} S(\mathbf{T})_{ave} &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{T}^T N_{il}(\mathbf{T} \mathbf{T}^T) \mathbf{T}] + \text{tr}[\mathbf{T}^T \mathbf{W}_o \mathbf{T}] \\ &+ \text{tr}[\mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T}] + \frac{L+1}{2} \end{aligned} \quad (35)$$

From (2) and (27), we can show that both transfer functions $\mathbf{H}(z)$ in (11) and $H_{il}(z)$ in (12a) are invariant under the coordinate transformation in (26).

In order to prevent overflow on the states, we impose a scaling condition that controls l_2 -norm of the impulse response from input to state, hence the energy flow in the signal path. Analytically this is performed by examining the input-to-state transfer function of the system as

$$(z\mathbf{I}_n - \hat{\mathbf{A}})^{-1} \hat{\mathbf{B}} = \sum_{k=0}^{\infty} \hat{\mathbf{A}}^k \hat{\mathbf{B}} z^{-(k+1)} \quad (36)$$

hence the input-to-state impulse response is given by $\{\hat{\mathbf{B}}, \hat{\mathbf{A}}\hat{\mathbf{B}}, \hat{\mathbf{A}}^2\hat{\mathbf{B}}, \dots\}$. The impulse response has n sequences (rows) of the form $\mathbf{e}_i^T \cdot \{\hat{\mathbf{B}}, \hat{\mathbf{A}}\hat{\mathbf{B}}, \hat{\mathbf{A}}^2\hat{\mathbf{B}}, \dots\}$ where $\mathbf{e}_i$ denotes the i th column of the identity matrix $\mathbf{I}_n$, and the l_2 -scaling requires that the l_2 -norm of each sequence be unity, namely

$$\mathbf{e}_i^T \left[\sum_{k=0}^{\infty} \hat{\mathbf{A}}^k \hat{\mathbf{B}} (\hat{\mathbf{A}}^k \hat{\mathbf{B}})^T \right] \mathbf{e}_i = 1 \quad \text{for } i = 1, 2, \dots, n \quad (37a)$$

From system model (2), it can readily be verified that

$$\sum_{k=0}^{\infty} \hat{\mathbf{A}}^k \hat{\mathbf{B}} (\hat{\mathbf{A}}^k \hat{\mathbf{B}})^T = \sum_{k=0}^{\infty} \mathbf{A}^k \mathbf{b} (\mathbf{A}^k \mathbf{b})^T = \mathbf{K}_c \quad (37b)$$

Therefore, the scaling conditions in (37a) can be reduced to

$$\mathbf{e}_i^T \mathbf{K}_c \mathbf{e}_i = 1 \quad \text{for } i = 1, 2, \dots, n \quad (37c)$$

Hence the l_2 -scaling on the vector impulse response $\overline{\mathbf{A}}^k \overline{\mathbf{b}}$ can be achieved by choosing a coordinate transformation in (26) so that

$$(\overline{\mathbf{K}}_c)_{ii} = (\mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T})_{ii} = 1 \quad \text{for } i = 1, 2, \dots, n \quad (38)$$

are satisfied.

Under these circumstances, the problem being considered here is to obtain an $n \times n$ coordinate transformation matrix $\mathbf{T}$ that minimizes the average l_2 -sensitivity measure $S(\mathbf{T})_{ave}$ in (35) subject to the l_2 -scaling constraints in (38).

Remark 1: The l_2 -sensitivity measure for the new realization $(\overline{\mathbf{A}}, \overline{\mathbf{b}}, \overline{\mathbf{c}}, d)_n$ characterized by (27) derived in [17], namely,

$$S_o(\mathbf{T}) = \text{tr}[\mathbf{T}^T \mathbf{N}_{11} (\mathbf{T} \mathbf{T}^T) \mathbf{T}] + \text{tr}[\mathbf{T}^T \mathbf{W}_o \mathbf{T}] + \text{tr}[\mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T}] + 1 \quad (39)$$

can be regarded as a special case of formula (35) with $L = 1$.

Remark 2: An average l_1/l_2 -sensitivity measure for the block-state realization in (2a) was proposed in [18],[19] as

$$M_{ave} = \text{tr}[\hat{\mathbf{W}}_o] \text{tr}[\hat{\mathbf{K}}_c] + \text{tr}[\hat{\mathbf{W}}_o] + \text{tr}[\hat{\mathbf{K}}_c] + \frac{L+1}{2} \quad (40a)$$

where

$$\begin{aligned} \hat{\mathbf{K}}_c &= \sum_{k=0}^{\infty} \hat{\mathbf{A}}^k \hat{\mathbf{B}} (\hat{\mathbf{A}}^k \hat{\mathbf{B}})^T = \sum_{l=1}^L \mathbf{K}_l \\ \hat{\mathbf{W}}_o &= \sum_{k=0}^{\infty} (\hat{\mathbf{C}} \hat{\mathbf{A}}^k)^T \hat{\mathbf{C}} \hat{\mathbf{A}}^k = \sum_{i=1}^L \mathbf{W}_i \end{aligned} \quad (40b)$$

with $\mathbf{K}_l$ and $\mathbf{W}_i$ obtained by solving the Lyapunov equations

$$\mathbf{K}_l = \hat{\mathbf{A}} \mathbf{K}_l \hat{\mathbf{A}}^T + \hat{\mathbf{b}}_l \hat{\mathbf{b}}_l^T, \quad \mathbf{W}_i = \hat{\mathbf{A}}^T \mathbf{W}_i \hat{\mathbf{A}} + \hat{\mathbf{c}}_i^T \hat{\mathbf{c}}_i \quad (40c)$$

for $i, l = 1, 2, \dots, L$.

3. Minimization of l_2 -Sensitivity Subject to l_2 -Scaling Constraints

A. Method 1: Using A Lagrange Function

We now consider the problem of minimizing the average l_2 -sensitivity measure in (35) subject to l_2 -scaling constraints in (38). Since the measure in (35) can be expressed in terms of matrix $\mathbf{P} = \mathbf{T} \mathbf{T}^T$ as

$$S(\mathbf{P})_{ave} = \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}(\mathbf{P}) \mathbf{P}] + \text{tr}[\mathbf{W}_o \mathbf{P}] + \text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] + \frac{L+1}{2} \quad (41)$$

the problem we deal with is a constrained nonlinear optimization problem where the variable is matrix $\mathbf{P}$.

It turns out the constraints in (38) is quite challenging. Here we handle this technical problem by a two-step procedure. First, by summing up the constraints in (38), we are led to a relaxed constraint as

$$\text{tr}[\mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T}] = \text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] = n \quad (42)$$

Consequently, the problem of minimizing the average l_2 -sensitivity measure in (35) subject to the constraints in (38) can be *relaxed* into the following problem:

$$\begin{aligned} &\text{minimize } S(\mathbf{P})_{ave} \text{ in (41) with respect to } \mathbf{P} \\ &\text{subject to } \text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] = n \end{aligned} \quad (43)$$

where $\mathbf{P} = \mathbf{T}^T$.

Note that although the solution $\mathbf{P}_{opt}$ to problem (43) satisfies constraint (42) which is equivalent to $\text{tr}[\mathbf{T}_{opt}^{-1} \mathbf{K}_c \mathbf{T}_{opt}^{-T}] = n$ where $\mathbf{T}_{opt} = \mathbf{P}_{opt}^{1/2}$. However, $\mathbf{T}_{opt}$ does not necessarily satisfy the l_2 -scaling constraints in (38). In the second step of our solution procedure detailed below, we will present a technique for obtaining an $n \times n$ orthogonal matrix $\mathbf{U}$ such that (38) is satisfied by $\mathbf{T} = \mathbf{P}_{opt}^{1/2} \mathbf{U}$ while the relation $\mathbf{T} \mathbf{T}^T = \mathbf{P}_{opt}$ is obviously maintained.

We now address problem (43) as the first step of our solution procedure. To this end, we define the Lagrange function of the problem as

$$\begin{aligned} J(\mathbf{P}, \lambda) &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{N}_{il}(\mathbf{P}) \mathbf{P}] + \text{tr}[\mathbf{W}_o \mathbf{P}] \\ &\quad + n + \frac{L+1}{2} + \lambda (\text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] - n) \end{aligned} \quad (44)$$

where λ is a Lagrange multiplier. It is well known that the solution of problem (43) must satisfy the Karush-Kuhn-Tucker (KKT) conditions $\partial J(\mathbf{P}, \lambda) / \partial \mathbf{P} = \mathbf{0}$ and $\partial J(\mathbf{P}, \lambda) / \partial \lambda = 0$, where

$$\begin{aligned} \frac{\partial J(\mathbf{P}, \lambda)}{\partial \mathbf{P}} &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \mathbf{N}_{il}(\mathbf{P}) + \mathbf{W}_o \\ &\quad - \mathbf{P}^{-1} \left[\frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \mathbf{M}_{il}(\mathbf{P}) + \lambda \mathbf{K}_c \right] \mathbf{P}^{-1} \quad (45) \\ \frac{\partial J(\mathbf{P}, \lambda)}{\partial \lambda} &= \text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] - n \end{aligned}$$

The matrices $\mathbf{M}_{il}(\mathbf{P})$ for $i, l = 1, 2, \dots, L$ in (45) or (31b) can be obtained by solving the Lyapunov equations

$$\begin{aligned} \mathbf{X}_{il} &= \begin{bmatrix} \mathbf{A}^L & \mathbf{A}^{L-l} \mathbf{b} \mathbf{c} \mathbf{A}^{i-1} \\ \mathbf{0} & \mathbf{A}^L \end{bmatrix} \mathbf{X}_{il} \begin{bmatrix} \mathbf{A}^L & \mathbf{A}^{L-l} \mathbf{b} \mathbf{c} \mathbf{A}^{i-1} \\ \mathbf{0} & \mathbf{A}^L \end{bmatrix}^T \\ &\quad + \begin{bmatrix} \mathbf{0} & \mathbf{0} \\ \mathbf{0} & \mathbf{P} \end{bmatrix} \end{aligned} \quad (46)$$

and then taking the upper-left $n \times n$ block of $\mathbf{X}_{il}$ as $\mathbf{M}_{il}(\mathbf{P})$:

$$\mathbf{M}_{il}(\mathbf{P}) = \begin{bmatrix} \mathbf{I}_n & \mathbf{0} \end{bmatrix} \mathbf{X}_{il} \begin{bmatrix} \mathbf{I}_n \\ \mathbf{0} \end{bmatrix} \quad (47)$$

From (45), the KKT conditions can be expressed compactly as

$$\mathbf{P} \mathbf{F}(\mathbf{P}) \mathbf{P} = \mathbf{G}(\mathbf{P}, \lambda), \quad \text{tr}[\mathbf{K}_c \mathbf{P}^{-1}] = n \quad (48a)$$

where

$$\begin{aligned} \mathbf{F}(\mathbf{P}) &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \mathbf{N}_{il}(\mathbf{P}) + \mathbf{W}_o \\ \mathbf{G}(\mathbf{P}, \lambda) &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \mathbf{M}_{il}(\mathbf{P}) + \lambda \mathbf{K}_c \end{aligned} \quad (48b)$$

Note that the first equation in (48a) is highly nonlinear with respect to $\mathbf{P}$. An effective approach to solving the first equation in (48a) is to *relax* it into the recursive second-order matrix equation

$$\mathbf{P}_{k+1} \mathbf{F}(\mathbf{P}_k) \mathbf{P}_{k+1} = \mathbf{G}(\mathbf{P}_k, \lambda_{k+1}) \quad (49)$$

where $\mathbf{P}_k$ is assumed to be known from the previous recursion. Note that if the matrix sequence $\{\mathbf{P}_k\}$ converges to its limit matrix, say $\mathbf{P}$, then (49) converges the first equation in (48a) as k goes to infinity. The solution $\mathbf{P}_{k+1}$ of (49) is then given by

$$\begin{aligned} \mathbf{P}_{k+1} &= \mathbf{F}(\mathbf{P}_k)^{-\frac{1}{2}} \left[\mathbf{F}(\mathbf{P}_k)^{\frac{1}{2}} \mathbf{G}(\mathbf{P}_k, \lambda_{k+1}) \mathbf{F}(\mathbf{P}_k)^{\frac{1}{2}} \right]^{\frac{1}{2}} \mathbf{F}(\mathbf{P}_k)^{-\frac{1}{2}} \\ &\quad (50) \end{aligned}$$

To derive a recursive formula for the Lagrange multiplier λ , we use (48a) to write

$$\text{tr}[\mathbf{P} \mathbf{F}(\mathbf{P})] = \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{M}_{il}(\mathbf{P}) \mathbf{P}^{-1}] + n\lambda \quad (51)$$

which naturally suggests the recursion for λ

$$\lambda_{k+1} = \frac{\text{tr}[\mathbf{P}_k \mathbf{F}(\mathbf{P}_k)] - \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L \text{tr}[\mathbf{M}_{il}(\mathbf{P}_k) \mathbf{P}_k^{-1}]}{n} \quad (52)$$

where $\mathbf{P}_0$ is an initial estimate such that $\mathbf{P}_0 = \mathbf{P}_0^T$.

Remark 3: If the initial estimate is chosen as $\mathbf{P}_0 = \mathbf{K}_c$, the related matrix $\mathbf{T}_0 = \mathbf{K}_c^{1/2} \mathbf{V}$ produces an input-normal realization $(\mathbf{T}_0^{-1} \mathbf{A} \mathbf{T}_0, \mathbf{T}_0^{-1} \mathbf{b}, \mathbf{c} \mathbf{T}_0, d)_n$ satisfying $\mathbf{T}_0^{-1} \mathbf{K}_c \mathbf{T}_0^{-T} = \mathbf{I}_n$ and $\mathbf{T}_0^T \mathbf{W}_o \mathbf{T}_0 = \mathbf{\Sigma}^2$ [21] where $\mathbf{V}$ and $\mathbf{\Sigma}$ are obtained by eigenvalue-eigenvector decomposition of $\mathbf{K}_c^{1/2} \mathbf{W}_o \mathbf{K}_c^{1/2}$ as

$$\mathbf{K}_c^{1/2} \mathbf{W}_o \mathbf{K}_c^{1/2} = \mathbf{V} \mathbf{\Sigma}^2 \mathbf{V}^T \quad (53)$$

with $n \times n$ diagonal $\mathbf{\Sigma}^2$ and orthogonal $\mathbf{V}$ composed of the eigenvalues and eigenvectors of $\mathbf{K}_c^{1/2} \mathbf{W}_o \mathbf{K}_c^{1/2}$, respectively.

The above iteration process continues until

$$|S(\mathbf{P}_{k+1})_{ave} - S(\mathbf{P}_k)_{ave}| + |n - \text{tr}[\mathbf{K}_c \mathbf{P}_{k+1}^{-1}]| < \varepsilon \quad (54)$$

is satisfied where $\varepsilon > 0$ is a prescribed tolerance. If the iteration is terminated at step k , then we take $\mathbf{P} = \mathbf{P}_k$ as the solution. Since $\mathbf{P} = \mathbf{T} \mathbf{T}^T$, the optimal $\mathbf{T}$ assumes the form

$$\mathbf{T} = \mathbf{P}^{\frac{1}{2}} \mathbf{U} \quad (55)$$

where $\mathbf{P}^{1/2}$ is the square root of matrix $\mathbf{P}$ obtained above, and $\mathbf{U}$ is an $n \times n$ orthogonal matrix.

As the second step of the solution procedure, we now turn our attention to the construction of the optimal coordinate transformation matrix $\mathbf{T}$ that solves the problem of minimizing the measure in (35) subject to the constraints in (38). To this end, we examine a procedure for determining the $n \times n$ orthogonal matrix $\mathbf{U}$ in (55) in order for the non-singular matrix $\mathbf{T}$ to satisfy the l_2 -scaling constraints in (38) without altering the optimal $\mathbf{P}$.

From (28) and (55), it follows that

$$\bar{\mathbf{K}}_c = \mathbf{T}^{-1} \mathbf{K}_c \mathbf{T}^{-T} = \mathbf{U}^T \mathbf{P}^{-\frac{1}{2}} \mathbf{K}_c \mathbf{P}^{-\frac{1}{2}} \mathbf{U} \quad (56)$$

In order to find an $n \times n$ orthogonal matrix $\mathbf{U}$ such that the matrix $\bar{\mathbf{K}}_c$ in (56) satisfies the scaling constraints in (38), we perform the eigenvalue-eigenvector decomposition for the positive-definite matrix $\mathbf{P}^{-1/2} \mathbf{K}_c \mathbf{P}^{-1/2}$ as

$$\mathbf{P}^{-\frac{1}{2}} \mathbf{K}_c \mathbf{P}^{-\frac{1}{2}} = \mathbf{R} \mathbf{\Theta} \mathbf{R}^T \quad (57)$$

where $\mathbf{\Theta} = \text{diag}\{\theta_1, \theta_2, \dots, \theta_n\}$ with $\theta_i > 0$ and $\mathbf{R}$ is an orthogonal matrix. Next, an orthogonal matrix $\mathbf{S}$ such that

$$(\mathbf{S} \mathbf{\Theta} \mathbf{S}^T)_{ii} = 1 \quad \text{for } i = 1, 2, \dots, n \quad (58)$$

can be obtained by numerical manipulations [3] (p.278). By using (56)-(58), it can be readily verified that the orthogonal matrix $\mathbf{U} = \mathbf{R} \mathbf{S}^T$ leads to a $\bar{\mathbf{K}}_c$ in (56) whose diagonal elements are equal to unity, hence the constraints in (38) are now satisfied. This matrix $\mathbf{T}$ together with (55) gives the solution to the problem of minimizing (35) subject to the constraints in (38) as

$$\mathbf{T} = \mathbf{P}^{\frac{1}{2}} \mathbf{R} \mathbf{S}^T \quad (59)$$

B. Method 2: Using A Quasi-Newton Algorithm

Since the state-space model in (1) is stable and controllable, the controllability Grammian $\mathbf{K}_c$ is positive-definite. Hence $\mathbf{K}_c^{1/2}$ satisfying $\mathbf{K}_c = \mathbf{K}_c^{1/2} \mathbf{K}_c^{1/2}$ is also positive-definite.

By defining

$$\hat{\mathbf{T}} = \mathbf{T}^T \mathbf{K}_c^{-\frac{1}{2}} \quad (60)$$

the l_2 -scaling constraints in (38) can be written as

$$(\hat{\mathbf{T}}^{-T} \hat{\mathbf{T}}^{-1})_{ii} = 1 \quad \text{for } i = 1, 2, \dots, n \quad (61)$$

The constraints in (61) simply state that each column in

matrix $\hat{\mathbf{T}}^{-1}$ must be a unit vector. If matrix $\hat{\mathbf{T}}^{-1}$ is assumed to have the form

$$\hat{\mathbf{T}}^{-1} = \left[\frac{\mathbf{t}_1}{\|\mathbf{t}_1\|}, \frac{\mathbf{t}_2}{\|\mathbf{t}_2\|}, \dots, \frac{\mathbf{t}_n}{\|\mathbf{t}_n\|} \right] \quad (62)$$

then (61) is always satisfied. Using (60), it follows from (35) that

$$S(\hat{\mathbf{T}})_{ave} = \text{tr}[\hat{\mathbf{T}}\hat{\mathbf{N}}(\hat{\mathbf{T}})\hat{\mathbf{T}}^T] + \text{tr}[\hat{\mathbf{T}}\tilde{\mathbf{W}}_o\hat{\mathbf{T}}^T] + n + \frac{L+1}{2} \quad (63a)$$

where

$$\begin{aligned} \text{tr}[\hat{\mathbf{T}}\hat{\mathbf{N}}(\hat{\mathbf{T}})\hat{\mathbf{T}}^T] &= \text{tr}[\hat{\mathbf{T}}^{-T}\hat{\mathbf{M}}(\hat{\mathbf{T}})\hat{\mathbf{T}}^{-1}] \\ [\hat{\mathbf{N}}(\hat{\mathbf{T}}), \hat{\mathbf{M}}(\hat{\mathbf{T}})] &= \frac{1}{L} \sum_{i=1}^L \sum_{l=1}^L [\hat{\mathbf{N}}_{il}(\hat{\mathbf{T}}), \hat{\mathbf{M}}_{il}(\hat{\mathbf{T}})] \\ \hat{\mathbf{N}}_{il}(\hat{\mathbf{T}}) &= \mathbf{K}_c^{\frac{1}{2}} \mathbf{N}_{il}(\mathbf{K}_c^{\frac{1}{2}} \hat{\mathbf{T}} \mathbf{K}_c^{\frac{1}{2}}) \mathbf{K}_c^{\frac{1}{2}} \\ \hat{\mathbf{M}}_{il}(\hat{\mathbf{T}}) &= \mathbf{K}_c^{-\frac{1}{2}} \mathbf{M}_{il}(\mathbf{K}_c^{\frac{1}{2}} \hat{\mathbf{T}} \mathbf{K}_c^{\frac{1}{2}}) \mathbf{K}_c^{-\frac{1}{2}} \\ \tilde{\mathbf{W}}_o &= \mathbf{K}_c^{\frac{1}{2}} \mathbf{W}_o \mathbf{K}_c^{\frac{1}{2}} \end{aligned} \quad (63b)$$

From the foregoing arguments, the problem of obtaining an $n \times n$ nonsingular matrix $\mathbf{T}$ which minimizes $S(\mathbf{T})_{ave}$ in (35) subject to the l_2 -scaling constraints in (38) can be converted into an unconstrained optimization problem of obtaining an $n \times n$ nonsingular matrix $\hat{\mathbf{T}}$ which minimizes $S(\hat{\mathbf{T}})_{ave}$ in (63a).

We now apply a quasi-Newton (known as the Broyden-Fletcher-Goldfarb-Shanno (BFGS)) algorithm [22] to minimize (63a) with respect to $\hat{\mathbf{T}}$ in (62). Let $\mathbf{x}$ be the column vector that collects the independent variables in matrix $\hat{\mathbf{T}}$, i.e.,

$$\mathbf{x} = (\mathbf{t}_1^T, \mathbf{t}_2^T, \dots, \mathbf{t}_n^T)^T \quad (64)$$

Then, $S(\hat{\mathbf{T}})_{ave}$ is a function of $\mathbf{x}$ and is denoted by $J_o(\mathbf{x})$. The algorithm starts with a trivial initial point $\mathbf{x}_0$ obtained from an initial assignment $\hat{\mathbf{T}} = \mathbf{I}_n$ (therefore $\mathbf{T} = \mathbf{K}_c^{1/2}$ in (60)). Then, in the k th iteration, a quasi-Newton algorithm updates the most recent point $\mathbf{x}_k$ to point $\mathbf{x}_{k+1}$ as [22]

$$\mathbf{x}_{k+1} = \mathbf{x}_k + \alpha_k \mathbf{d}_k \quad (65a)$$

where

$$\begin{aligned} \mathbf{d}_k &= -\mathbf{S}_k \nabla J_o(\mathbf{x}_k), \quad \alpha_k = \arg[\min_{\alpha} J_o(\mathbf{x}_k + \alpha \mathbf{d}_k)] \\ \mathbf{S}_{k+1} &= \mathbf{S}_k + \left(1 + \frac{\gamma_k^T \mathbf{S}_k \gamma_k}{\gamma_k^T \delta_k} \right) \frac{\delta_k \delta_k^T}{\gamma_k^T \delta_k} - \frac{\delta_k \gamma_k^T \mathbf{S}_k + \mathbf{S}_k \gamma_k \delta_k^T}{\gamma_k^T \delta_k} \\ \mathbf{S}_0 &= \mathbf{I}, \quad \delta_k = \mathbf{x}_{k+1} - \mathbf{x}_k, \quad \gamma_k = \nabla J_o(\mathbf{x}_{k+1}) - \nabla J_o(\mathbf{x}_k) \end{aligned} \quad (65b)$$

In the above, $\nabla J_o(\mathbf{x})$ is the gradient of $J_o(\mathbf{x})$ with respect to $\mathbf{x}$, and $\mathbf{S}_k$ is a positive-definite approximation of the inverse Hessian matrix of $J_o(\mathbf{x})$. Outlined below is an algorithm to compute an approximate local minimizer α_k for $J_o(\mathbf{x}_k +$

$\alpha \mathbf{d}_k$). In Section 4, $\alpha_{init} = 2^{-25}$ and $N = 30$ were chosen.

Algorithm 1 Search $\alpha_k = \arg[\min_{\alpha} J_o(\mathbf{x}_k + \alpha \mathbf{d}_k)]$

```

1:  $\alpha_k \leftarrow 0$ 
2:  $\alpha \leftarrow \alpha_{init}$   $\triangleright \alpha_{init}$  is a small positive value
3:  $J_{o,min} \leftarrow J_o(\mathbf{x}_k)$ 
4: for  $n = 1 \dots N$  do
5:   if  $J_{o,min} > J_o(\mathbf{x}_k + \alpha \mathbf{d}_k)$  then
6:      $J_{o,min} \leftarrow J_o(\mathbf{x}_k + \alpha \mathbf{d}_k)$ 
7:      $\alpha_k \leftarrow \alpha$ 
8:   end if
9:    $\alpha \leftarrow 2\alpha$ 
10: end for
11:  $\Delta\alpha \leftarrow \alpha_k/2$ 
12: for  $n = 1 \dots N$  do
13:   if  $J_{o,min} > J_o(\mathbf{x}_k + (\alpha_k + \Delta\alpha) \mathbf{d}_k)$  then
14:      $J_{o,min} \leftarrow J_o(\mathbf{x}_k + (\alpha_k + \Delta\alpha) \mathbf{d}_k)$ 
15:      $\alpha_k \leftarrow \alpha_k + \Delta\alpha$ 
16:   else if  $J_{o,min} > J_o(\mathbf{x}_k + (\alpha_k - \Delta\alpha) \mathbf{d}_k)$  then
17:      $J_{o,min} \leftarrow J_o(\mathbf{x}_k + (\alpha_k - \Delta\alpha) \mathbf{d}_k)$ 
18:      $\alpha_k \leftarrow \alpha_k - \Delta\alpha$ 
19:   end if
20:    $\Delta\alpha \leftarrow \Delta\alpha/2$ 
21: end for

```

This iteration process continues until

$$|J_o(\mathbf{x}_{k+1}) - J_o(\mathbf{x}_k)| < \varepsilon \quad (66)$$

is satisfied where $\varepsilon > 0$ is a prescribed tolerance. If the iteration is terminated at step k , then $\mathbf{x}_k$ is taken to be the solution to the minimization problem.

The implementation of (65a) requires the computation of $\nabla J_o(\mathbf{x})$ in each iteration, and hence the availability of a closed-form formulation to compute $\nabla J_o(\mathbf{x})$ will make the algorithm considerably faster relative to the evaluation of $\nabla J_o(\mathbf{x})$ based on numerical differentiation.

Each term of the objective function in (63a) has the form $J(\mathbf{x}) = \text{tr}[\hat{\mathbf{T}}\mathbf{N}\hat{\mathbf{T}}^T]$ (or $J(\mathbf{x}) = \text{tr}[\hat{\mathbf{T}}^{-T}\mathbf{M}\hat{\mathbf{T}}^{-1}]$) which, in the light of (62), can be expressed as

$$J(\mathbf{x}) = \text{tr} \left[\left[\frac{\mathbf{t}_1}{\|\mathbf{t}_1\|}, \frac{\mathbf{t}_2}{\|\mathbf{t}_2\|}, \dots, \frac{\mathbf{t}_n}{\|\mathbf{t}_n\|} \right]^{-1} \cdot N \left[\frac{\mathbf{t}_1}{\|\mathbf{t}_1\|}, \frac{\mathbf{t}_2}{\|\mathbf{t}_2\|}, \dots, \frac{\mathbf{t}_n}{\|\mathbf{t}_n\|} \right]^{-T} \right] \quad (67)$$

To compute $\partial J(\mathbf{x})/\partial t_{pq}$, we perturb the p th component of vector $\mathbf{t}_q$ by a small amount, say Δ , and keep the rest of $\hat{\mathbf{T}}$ unchanged. If we denote the perturbed q th column of $\hat{\mathbf{T}}^{-1}$ by $\tilde{\mathbf{t}}_q/\|\tilde{\mathbf{t}}_q\|$, then a linear approximation of $\tilde{\mathbf{t}}_q/\|\tilde{\mathbf{t}}_q\|$ can be obtained as

$$\frac{\tilde{\mathbf{t}}_q}{\|\tilde{\mathbf{t}}_q\|} \simeq \frac{\mathbf{t}_q}{\|\mathbf{t}_q\|} + \Delta \partial \left\{ \frac{\mathbf{t}_q}{\|\mathbf{t}_q\|} \right\} / \partial t_{pq} = \frac{\mathbf{t}_q}{\|\mathbf{t}_q\|} - \Delta \mathbf{g}_{pq} \quad (68a)$$

where

$$\mathbf{g}_{pq} = -\partial \left\{ \frac{\mathbf{t}_q}{\|\mathbf{t}_q\|} \right\} / \partial t_{pq} = \frac{1}{\|\mathbf{t}_q\|^3} (t_{pq} \mathbf{t}_q - \|\mathbf{t}_q\|^2 \mathbf{e}_p) \quad (68b)$$

Now let $\hat{T}_{pq}$ be the matrix obtained from $\hat{T}$ with a perturbed (p, q) th component, then we obtain

$$\hat{T}_{pq}^{-1} = \hat{T}^{-1} - \Delta g_{pq} e_q^T \quad (69)$$

and the matrix inversion formula [23] (p. 655) gives

$$\hat{T}_{pq} = \hat{T} + \frac{\Delta \hat{T} g_{pq} e_q^T \hat{T}}{1 - \Delta e_q^T \hat{T} g_{pq}} \simeq \hat{T} + \Delta \hat{T} g_{pq} e_q^T \hat{T} \quad (70)$$

For convenience, we define $\hat{T}_{pq} = \hat{T} + \Delta S$ and write

$$\hat{T}_{pq} N \hat{T}_{pq}^T = \hat{T} N \hat{T}^T + \Delta S N \hat{T}^T + \Delta \hat{T} N S^T + \Delta^2 S N S^T \quad (71)$$

which implies that

$$\text{tr}[\hat{T}_{pq} N \hat{T}_{pq}^T] - \text{tr}[\hat{T} N \hat{T}^T] \simeq \Delta \text{tr}[S(N + N^T) \hat{T}^T] \quad (72)$$

provided that Δ is sufficiently small. Hence

$$\begin{aligned} \frac{\partial \text{tr}[\hat{T} N \hat{T}^T]}{\partial t_{pq}} &= \lim_{\Delta \rightarrow 0} \frac{\text{tr}[\hat{T}_{pq} N \hat{T}_{pq}^T] - \text{tr}[\hat{T} N \hat{T}^T]}{\Delta} \\ &= \text{tr}[S(N + N^T) \hat{T}^T] \end{aligned} \quad (73)$$

Similarly, if we define $\hat{T}_{pq}^{-1} = \hat{T}^{-1} - \Delta S$, then we arrive at

$$\begin{aligned} \frac{\partial \text{tr}[\hat{T}^{-T} M \hat{T}^{-1}]}{\partial t_{pq}} &= \lim_{\Delta \rightarrow 0} \frac{\text{tr}[\hat{T}_{pq}^{-T} M \hat{T}_{pq}^{-1}] - \text{tr}[\hat{T}^{-T} M \hat{T}^{-1}]}{\Delta} \\ &= -\text{tr}[S^T (M + M^T) \hat{T}^{-1}] \end{aligned} \quad (74)$$

Referring to (73), it follows from (63a) and (70) that

$$\begin{aligned} \frac{\partial \text{tr}[\hat{T} \hat{N}(\hat{T}) \hat{T}^T]}{\partial t_{pq}} &= 2 \text{tr}[\hat{T} g_{pq} e_q^T \hat{N}(\hat{T}) \hat{T}^T] \\ &= 2 e_q^T \hat{N}(\hat{T}) \hat{T}^T \hat{T} g_{pq} \\ \frac{\partial \text{tr}[\hat{T} \tilde{W}_o \hat{T}^T]}{\partial t_{pq}} &= 2 \text{tr}[\hat{T} g_{pq} e_q^T \tilde{W}_o \hat{T}^T] \\ &= 2 e_q^T \hat{T} \tilde{W}_o \hat{T}^T \hat{T} g_{pq} \end{aligned} \quad (75)$$

Referring to (74), it follows from (63a) and (69) that

$$\begin{aligned} \frac{\partial \text{tr}[\hat{T}^{-T} \hat{M}(\hat{T}) \hat{T}^{-1}]}{\partial t_{pq}} &= -2 \text{tr}[e_q g_{pq}^T \hat{M}(\hat{T}) \hat{T}^{-1}] \\ &= -2 g_{pq}^T \hat{M}(\hat{T}) \hat{T}^{-1} e_q = -2 e_q^T \hat{T}^{-T} \hat{M}(\hat{T}) g_{pq} \end{aligned} \quad (76)$$

The gradient of $J_o(\mathbf{x})$ with respect to t_{pq} can be evaluated using closed-form expressions as

$$\begin{aligned} \frac{\partial J_o(\mathbf{x})}{\partial t_{pq}} &= \lim_{\Delta \rightarrow 0} \frac{S(\hat{T}_{pq})_{ave} - S(\hat{T})_{ave}}{\Delta} \\ &= 2(\beta_1 - \beta_2 + \beta_3) \end{aligned} \quad (77a)$$

where

$$\begin{aligned} \beta_1 &= e_q^T \hat{T} \hat{N}(\hat{T}) \hat{T}^T \hat{T} g_{pq} \\ \beta_2 &= e_q^T \hat{T}^{-T} \hat{M}(\hat{T}) g_{pq}, \quad \beta_3 = e_q^T \hat{T} \tilde{W}_o \hat{T}^T \hat{T} g_{pq} \end{aligned} \quad (77b)$$

4. A Case Study

Consider the 4th-order Butterworth low-pass digital filter $(A_o, b_o, c_o, d)_4$ with a narrow normalized passband of 0.05, described by

$$\begin{aligned} A_o &= \begin{bmatrix} 3.589734 & 1 & 0 & 0 \\ -4.851276 & 0 & 1 & 0 \\ 2.924053 & 0 & 0 & 1 \\ -0.663010 & 0 & 0 & 0 \end{bmatrix}, \quad b_o = 10^{-3} \begin{bmatrix} 0.237096 \\ 0.035885 \\ 0.216300 \\ 0.010527 \end{bmatrix} \\ c_o &= \begin{bmatrix} 1 & 0 & 0 & 0 \end{bmatrix}, \quad d = 3.123898 \times 10^{-5} \end{aligned}$$

whose numerator b and denominator a were found using MATLAB as $[b, a] = \text{butter}(4, 0.05)$.

Applying a coordinate transformation defined by

$$T_s = \text{diag}\{0.226458, 0.588059, 0.513017, 0.150144\}$$

to the above state-space model, we obtained an equivalent state-space model $(A, b, c, d)_4$ that satisfies l_2 -scaling constraints as

$$\begin{aligned} A &= \begin{bmatrix} 3.589734 & 2.596768 & 0 & 0 \\ -1.868197 & 0 & 0.872390 & 0 \\ 1.290747 & 0 & 0 & 0.292669 \\ -1.000000 & 0 & 0 & 0 \end{bmatrix} \\ b &= 10^{-3} \begin{bmatrix} 1.046973 & 0.061023 & 0.421624 & 0.070114 \end{bmatrix}^T \\ c &= \begin{bmatrix} 0.226458 & 0 & 0 & 0 \end{bmatrix}, \quad d = 3.123898 \times 10^{-5} \end{aligned}$$

where $A = T_s^{-1} A_o T_s$, $b = T_s^{-1} b_o$, $c = c_o T_s$ and the eigenvalues of A have the magnitudes 0.941824 and 0.864550. By solving the Lyapunov equations in (22c), the controllability and observability Grammians K_c and W_o were computed as

$$\begin{aligned} K_c &= \begin{bmatrix} 1.000000 & -0.999248 & 0.997433 & -0.994918 \\ -0.999248 & 1.000000 & -0.999452 & 0.998047 \\ 0.997433 & -0.999452 & 1.000000 & -0.999566 \\ -0.994918 & 0.998047 & -0.999566 & 1.000000 \end{bmatrix} \\ W_o &= 10^4 \begin{bmatrix} 1.063597 & 2.747677 & 2.360090 & 0.672994 \\ 2.747677 & 7.172055 & 6.224575 & 1.793652 \\ 2.360090 & 6.224575 & 5.458399 & 1.589267 \\ 0.672994 & 1.793652 & 1.589267 & 0.467539 \end{bmatrix} \end{aligned}$$

Since $\sqrt{2}n = 5.656854$ for the filter order $n = 4$, the blocklength is set to $L = 6$ in this section. The block-state realization was then constructed from (2b) as

$$\hat{\mathbf{A}} = \begin{bmatrix} 42.27560 & 82.60016 & 50.68938 & 9.51583 \\ -37.89405 & -71.90962 & -42.32252 & -7.44160 \\ 34.50854 & 64.08980 & 36.64013 & 6.17955 \\ -31.80883 & -58.10401 & -32.51401 & -5.32723 \end{bmatrix}$$

$$\hat{\mathbf{B}} = 10^{-3} \begin{bmatrix} 50.9309 & 33.5423 & 19.7751 & \\ -42.9331 & -26.7551 & -14.4198 & \\ 37.5070 & 22.6166 & 11.6789 & \\ -33.5423 & -19.7751 & -9.93632 & \\ & 9.93632 & 3.91682 & 1.04697 \\ & -6.12056 & -1.58813 & 0.06102 \\ & 4.74920 & 1.37190 & 0.42162 \\ & -3.91682 & -1.04697 & 0.07011 \end{bmatrix}$$

$$\hat{\mathbf{C}} = \begin{bmatrix} 0.226458 & 0 & 0 & 0 \\ 0.812924 & 0.588059 & 0 & 0 \\ 1.819571 & 2.110976 & 0.513017 & 0 \\ 3.250232 & 4.725005 & 1.841595 & 0.150144 \\ 5.067114 & 8.440099 & 4.122049 & 0.538977 \\ 7.203366 & 13.158123 & 7.363061 & 1.206395 \end{bmatrix}$$

$$\begin{bmatrix} \hat{d}_{11} & \hat{d}_{21} & \hat{d}_{31} \\ \hat{d}_{41} & \hat{d}_{51} & \hat{d}_{61} \end{bmatrix} = 10^{-3} \begin{bmatrix} 0.031239 & 0.237096 & 0.886995 \\ 2.250160 & 4.478225 & 7.595914 \end{bmatrix}$$

where the eigenvalues of $\hat{\mathbf{A}}$ have the magnitudes 0.697942 and 0.417582.

The average l_2 -sensitivity measure for the block-state realization described by (11) was then computed from (25) as

$$S_{ave} = 40.933372 \times 10^4 \quad (27.330613 \times 10^4)$$

which was compared with the l_2 -sensitivity of the state-space model, shown in (23) as

$$S_o = 977.917589 \times 10^4 \quad (488.229633 \times 10^4)$$

where the values in the round brackets are the sensitivities that have taken the trivial elements with 0 and ± 1 values in the coefficient matrices into account [24].

A) The Use of A Lagrange Function

The recursive matrix equation in (50) together with (52) was applied to minimize (44) with tolerance $\varepsilon = 10^{-8}$ in (54) and initial estimate $\mathbf{P}_0 = \mathbf{K}_c$. It took the algorithm 76 iterations to converge to the solution

$$\mathbf{P} = \begin{bmatrix} 0.875338 & -0.904805 & 0.932402 & -0.958480 \\ -0.904805 & 0.938156 & -0.969673 & 0.999738 \\ 0.932402 & -0.969673 & 1.005245 & -1.039528 \\ -0.958480 & 0.999738 & -1.039528 & 1.078311 \end{bmatrix}$$

and from (59)

$$\mathbf{T} = \begin{bmatrix} 0.327836 & 0.559715 & 0.663273 & 0.121038 \\ -0.313941 & -0.586064 & -0.699777 & -0.080238 \\ 0.301577 & 0.604055 & 0.740336 & 0.036271 \\ -0.294198 & -0.614717 & -0.783389 & 0.013515 \end{bmatrix}$$

was obtained. A new block-state realization was obtained from (27) and (30) as

$$\hat{\mathbf{A}}_{new} = \begin{bmatrix} 0.397376 & -0.473485 & 0.232855 & 0.220533 \\ 0.537510 & 0.388153 & -0.130919 & 0.201674 \\ -0.028026 & 0.015221 & 0.446037 & 0.327314 \\ 0.021347 & -0.406390 & 0.101068 & 0.447310 \end{bmatrix}$$

$$\hat{\mathbf{B}}_{new} = \begin{bmatrix} 0.040743 & -0.020980 & -0.090083 \\ -0.268273 & -0.280345 & -0.281986 \\ 0.241915 & 0.256942 & 0.271631 \\ 0.225341 & 0.222337 & 0.222858 \\ & -0.166178 & -0.248790 & -0.337393 \\ & -0.271737 & -0.248370 & -0.210986 \\ & 0.284592 & 0.293878 & 0.296858 \\ & 0.229256 & 0.244342 & 0.271408 \end{bmatrix}$$

$$\hat{\mathbf{C}}_{new} = \begin{bmatrix} 0.074241 & 0.126752 & 0.150203 & 0.027410 \\ 0.081890 & 0.110365 & 0.127680 & 0.051210 \\ 0.088513 & 0.091165 & 0.109465 & 0.069463 \\ 0.093382 & 0.070177 & 0.095119 & 0.083100 \\ 0.096040 & 0.048327 & 0.084164 & 0.092889 \\ 0.096263 & 0.026433 & 0.076108 & 0.099465 \end{bmatrix}$$

and the new controllability Grammian was found to be

$$\bar{\mathbf{K}}_c = \begin{bmatrix} 1.000000 & 0.300486 & 0.342142 & 0.474457 \\ 0.300486 & 1.000000 & -0.342142 & -0.474457 \\ 0.342142 & -0.342142 & 1.000000 & 0.910932 \\ 0.474457 & -0.474457 & 0.910932 & 1.000000 \end{bmatrix}$$

The average l_2 -sensitivity for the block-state realization described by (11) was then computed from (35) as

$$S(\mathbf{T})_{ave} = 8.759262$$

which was compared with the l_2 -sensitivity of a new state-space model $(\bar{\mathbf{A}}, \bar{\mathbf{b}}, \bar{\mathbf{c}}, d)_n$, shown in (39) as

$$S_o(\mathbf{T}) = 29.500327$$

The profiles of $S(\mathbf{P})_{ave}$ in (41) and $\text{tr}[\mathbf{K}_c \mathbf{P}^{-1}]$ during the first 76 iterations of the algorithm are depicted in Fig. 2 and Fig. 3, respectively.

B) The Use of A Quasi-Newton Algorithm

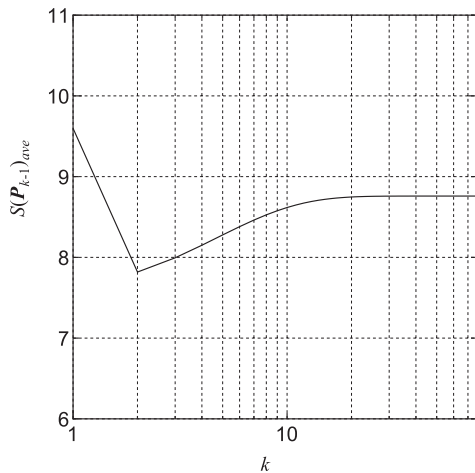
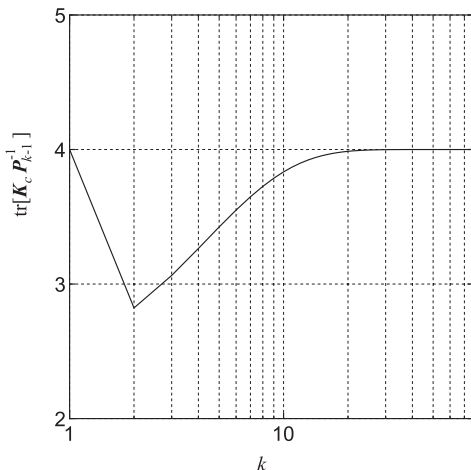
The quasi-Newton algorithm was applied to minimize (63a) with initial assignment $\hat{\mathbf{T}} = \mathbf{I}_4$, and tolerance $\varepsilon = 10^{-8}$ in (66). It took the algorithm 22 iterations to converge to the solution

$$\hat{\mathbf{T}} = \begin{bmatrix} 0.571134 & -0.415538 & -1.326551 & 0.389172 \\ 0.407792 & 0.933965 & -1.087187 & -2.120849 \\ 0.043983 & -0.117207 & 3.040058 & 2.480486 \\ -0.373909 & 0.428710 & 1.758732 & 2.561697 \end{bmatrix}$$

which is equivalent to

$$\mathbf{T} = \begin{bmatrix} -0.307934 & 0.205102 & 0.403183 & -0.758913 \\ 0.350071 & -0.233366 & -0.379012 & 0.785847 \\ -0.385730 & 0.267285 & 0.350018 & -0.813980 \\ 0.414649 & -0.303751 & -0.313832 & 0.845979 \end{bmatrix}$$

where $\mathbf{P} = \mathbf{T}\mathbf{T}^T$ was computed as

Fig. 2 Profile of $S(\mathbf{P})_{ave}$ during the first 76 iterations.Fig. 3 Profile of $\text{tr}[\mathbf{K}_c \mathbf{P}^{-1}]$ during the first 76 iterations.

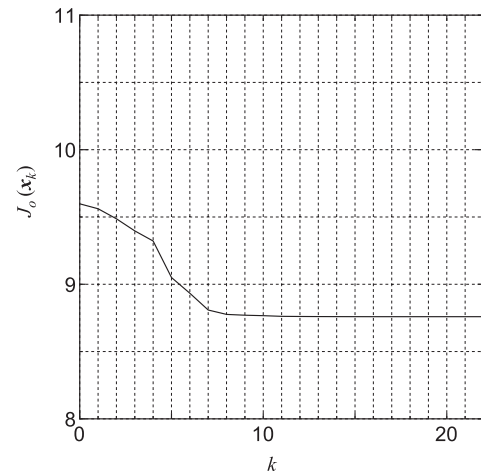
$$\mathbf{P} = \begin{bmatrix} 0.875395 & -0.904863 & 0.932461 & -0.958540 \\ -0.904863 & 0.938216 & -0.969734 & 0.999798 \\ 0.932461 & -0.969734 & 1.005306 & -1.039588 \\ -0.958540 & 0.999798 & -1.039588 & 1.078370 \end{bmatrix}$$

which is practically identical to matrix $\mathbf{P}$ in Section IV-A.

A new block-state realization was then obtained from (27) and (30) as

$$\hat{\mathbf{A}}_{new} = \begin{bmatrix} 0.426110 & 0.516400 & -0.130405 & 0.290374 \\ -0.269864 & 0.389027 & 0.160428 & 0.110379 \\ 0.450255 & 0.105445 & 0.449322 & -0.002678 \\ -0.390308 & -0.117906 & -0.316171 & 0.414417 \end{bmatrix}$$

$$\hat{\mathbf{B}}_{new} = \begin{bmatrix} 0.285271 & 0.262568 & 0.228614 \\ 0.268791 & 0.312840 & 0.356168 \\ 0.170038 & 0.164974 & 0.166233 \\ -0.019883 & 0.021455 & 0.065752 \\ 0.183139 & 0.126272 & 0.058614 \\ 0.396886 & 0.432672 & 0.460722 \\ 0.176240 & 0.197682 & 0.233484 \\ 0.113488 & 0.165557 & 0.223393 \end{bmatrix}$$

Fig. 4 Profile of $J_o(\mathbf{x})$ during the first 22 iterations.

$$\hat{\mathbf{C}}_{new} = \begin{bmatrix} -0.069734 & 0.046447 & 0.091304 & -0.171862 \\ -0.044464 & 0.029499 & 0.104875 & -0.154814 \\ -0.019201 & 0.017688 & 0.113099 & -0.139578 \\ 0.005132 & 0.010595 & 0.117075 & -0.125515 \\ 0.027790 & 0.007686 & 0.117717 & -0.112175 \\ 0.048199 & 0.008352 & 0.115785 & -0.099257 \end{bmatrix}$$

Moreover, the new controllability Grammian was found to be

$$\bar{\mathbf{K}}_c = \begin{bmatrix} 1.000000 & 0.609669 & 0.656806 & -0.159006 \\ 0.609669 & 1.000000 & 0.530787 & 0.197938 \\ 0.656806 & 0.530787 & 1.000000 & -0.674646 \\ -0.159006 & 0.197938 & -0.674646 & 1.000000 \end{bmatrix}$$

and the average l_2 -sensitivity measure for the block-state realization described by (11) was computed from (35) as

$$S(\mathbf{T})_{ave} = 8.759262$$

which was compared with the l_2 -sensitivity of a new state-space model $(\bar{\mathbf{A}}, \bar{\mathbf{b}}, \bar{\mathbf{c}}, d)_n$, shown in (39) as

$$S_o(\mathbf{T}) = 29.500303$$

The profile of the average l_2 -sensitivity measure $J_o(\mathbf{x})$ during the first 22 iterations of the algorithm is depicted in Fig. 4.

We see that as far as the digital filter examined in this example is concerned, the two algorithms offer practically the same performance in terms of l_2 -sensitivity minimization. Concerning the computational complexity, the 76 iterations required by the constrained algorithm consumed 4.09375 seconds of CPU time versus 41.12500 seconds for the unconstrained algorithm to finish its 22 iterations on a PC with a 3.3GHz i5-3550 CPU and 8GB memory. Obviously, the constrained method is considerably more efficient than its unconstrained counterpart. However, for digital filter of large scales, it is quite likely that the algorithms' complexity will change due to the fact that the constrained algorithm heavily involves expensive matrix manipulations such as computing square root as well as inverse of matrices. For this reason,

we believe that together the two algorithms can serve for l_2 -sensitivity minimization for a wide range of filter scales.

5. Conclusion

The l_2 -sensitivity in the block-state realization derived from a given SISO state-space model have been analyzed. The problem of minimizing the average l_2 -sensitivity measure subject to the l_2 -norm dynamic-range scaling constraints has then been solved by two techniques. One has been based on a Lagrange function, while the other has relied on an efficient quasi-Newton algorithm. Finally, a numerical example has been presented to demonstrate the validity and effectiveness of the proposed techniques.

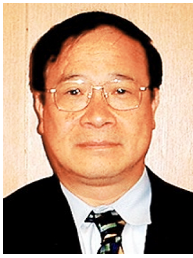
References

- [1] S.Y. Hwang, "Roundoff noise in state-space digital filtering: A general analysis," *IEEE Trans. Acoust., Speech, Signal Process.*, vol.ASSP-24, no.3, pp.256–262, June 1976.
- [2] C.T. Mullis and R.A. Roberts, "Synthesis of minimum roundoff noise fixed point digital filters," *IEEE Trans. Circuits Syst.*, vol.CAS-23, no.9, pp.551–562, Sept. 1976.
- [3] S.Y. Hwang, "Minimum uncorrelated unit noise in state-space digital filtering," *IEEE Trans. Acoust., Speech, Signal Process.*, vol.ASSP-25, no.4, pp.273–281, Aug. 1977.
- [4] C.W. Barnes and S. Shinnaka, "Finite word effects in block-state realization of fixed-point digital filters," *IEEE Trans. Circuits Syst.*, vol.CAS-27, no.5, pp.345–349, May 1980.
- [5] C.W. Barnes and S. Shinnaka, "Block-shift invariance and block implementation of discrete-time filters," *IEEE Trans. Circuits Syst.*, vol.CAS-27, no.8, pp.667–672, Aug. 1980.
- [6] J. Zeman and A.G. Lindgren, "Fast digital filters with low roundoff noise," *IEEE Trans. Circuits Syst.*, vol.CAS-28, no.7, pp.716–723, July 1981.
- [7] L. Thiele, "Design of sensitivity and round-off noise optimal state-space discrete systems," *Int. J. Circuit Theory Appl.*, vol.12, no.1, pp.39–46, Jan. 1984.
- [8] V. Tsvanoglu and L. Thiele, "Optimal design of state-space digital filters by simultaneous minimization of sensitivity and roundoff noise," *IEEE Trans. Circuits Syst.*, vol.CAS-31, no.10, pp.884–888, Oct. 1984.
- [9] L. Thiele, "On the sensitivity of linear state-space systems," *IEEE Trans. Circuits Syst.*, vol.CAS-33, no.5, pp.502–510, May 1986.
- [10] M. Iwatsuki, M. Kawamata, and T. Higuchi, "Statistical sensitivity and minimum sensitivity structures with fewer coefficients in discrete time linear systems," *IEEE Trans. Circuits Syst.*, vol.CAS-37, no.1, pp.72–80, Jan. 1989.
- [11] G. Li, B.D.O. Anderson, M. Gevers, and J.E. Perkins, "Optimal FWL design of state-space digital systems with weighted sensitivity minimization and sparseness consideration," *IEEE Trans. Circuits Syst. I, Fundam. Theory Appl.*, vol.39, no.5, pp.365–377, May 1992.
- [12] W.-Y. Yan and J.B. Moore, "On L^2 -sensitivity minimization of linear state-space systems," *IEEE Trans. Circuits Syst. I, Fundam. Theory Appl.*, vol.39, no.8, pp.641–648, Aug. 1992.
- [13] G. Li and M. Gevers, "Optimal synthetic FWL design of state-space digital filters," *Proc. ICASSP 1992*, vol.4, pp.429–432, 1992.
- [14] M. Gevers and G. Li, *Parameterizations in Control, Estimation and Filtering Problems: Accuracy Aspects*, Springer-Verlag, New York, 1993.
- [15] T. Hinamoto, S. Yokoyama, T. Inoue, W. Zeng and W.-S. Lu, "Analysis and minimization of L_2 -sensitivity for linear systems and two-dimensional state-space filters using general controllability and observability Gramians," *IEEE Trans. Circuits Syst. I, Fundam. Theory Appl.*, vol.49, no.9, pp.1279–1289, Sept. 2002.
- [16] T. Hinamoto, H. Ohnishi, and W.-S. Lu, "Minimization of L_2 -sensitivity for state-space digital filters subject to L_2 -dynamic-range scaling constraints," *IEEE Trans. Circuits Syst.-II, Exp. Briefs*, vol.52, no.10, pp.641–645, Oct. 2005.
- [17] T. Hinamoto, K. Iwata, and W.-S. Lu, " L_2 -sensitivity minimization of one- and two-dimensional state-space digital filters subject to L_2 -scaling constraints," *IEEE Trans. Signal Process.*, vol.54, no.5, pp.1804–1812, May 2006.
- [18] Y. Tsunekawa, J. Shida, M. Kawamata, and T. Higuchi, "Coefficient sensitivity analysis for block-state realizations of state-space digital filters," *IEICE Trans. Fundamentals (Japanese edition)*, vol.J74-A, no.7, pp.996–1005, July 1991, also *Electronics and Communications in Japan*, part 3, vol.75, no.3, pp.13–25, 1992.
- [19] Y. Tsunekawa, J. Shida, M. Kawamata, and T. Higuchi, "Minimization of coefficient sensitivity for block-state realizations of state-space digital filters," *IEICE Trans. Fundamentals (Japanese edition)*, vol.J75-A, no.3, pp.506–515, March 1992, also *Electronics and Communications in Japan*, part 3, vol.75, no.11, pp.88–101, 1992.
- [20] T. Hinamoto and W.-S. Lu, *Digital Filter Design and Realization*, River Publishers, Denmark/The Netherlands, 2017.
- [21] L.M. Silverman, "Optimal approximation of linear systems," *Proc. Joint Automat. Contr. Conf.*, S.F., FA8-A, 1980.
- [22] R. Fletcher, *Practical Methods of Optimization*, 2nd ed., Wiley, New York, 1987.
- [23] T. Kailath, *Linear System*, Prentice-Hall, Englewood Cliffs, N.J., 1980.
- [24] C. Xiao, "Improved L_2 -sensitivity for state-space digital systems," *IEEE Trans. Signal Process.*, vol.45, no.4, pp.837–840, April 1997.



Akimitsu Doi received the B.E. degree in Electronic Engineering from Tottori University in 1992, and M.E. and Ph.D. degrees in System Engineering from Hiroshima University in 1994 and 1997, respectively. From 1997 to 2002, he was a Research Associate with the Faculty of Engineering, Hiroshima University, Hiroshima, Japan. From 2002 to 2005, he was a Lecturer in the Faculty of Engineering, Hiroshima Institute of Technology, Hiroshima, Japan. From 2005 to 2015, he was an Associate Professor and since

October 2015, he has been a Professor at the same Institute. His research interests are in the area of digital signal and image processing.



Takao Hinamoto received the B.E. degree from Okayama University, Okayama, Japan, in 1969, the M.E. degree from Kobe University, Kobe, Japan, in 1971, and the Dr. Eng. degree from Osaka University, Osaka, Japan, in 1977, all in electrical engineering. From 1972 to 1988, he was with the Faculty of Engineering, Kobe University, Kobe, Japan. From 1979 to 1981, he was a Visiting Member of Staff in the Department of Electrical Engineering, Queen's University, Kingston, ON, Canada, on leave from Kobe

University. During 1988–1991, he was Professor of Electronic Circuits in the Faculty of Engineering, Tottori University, Tottori, Japan. During 1992–2009, he was Professor of Electronic Control in the Department of Electrical Engineering, Hiroshima University, Hiroshima, Japan. Since 2009, he has been a Professor Emeritus of Hiroshima University. His research interests include digital signal processing, system theory, and control engineering. He has published over 400 papers in these areas. He is co-editor and co-author of *Two-Dimensional Signal and Image Processing* (Tokyo, Japan: SICE, 1996). Dr. Hinamoto is a Life Fellow of the IEEE and a Fellow of the Society of Instrument and Control Engineers (SICE).



Wu-Sheng Lu received his undergraduate degree in mathematics from Fudan University, Shanghai, China, in 1964, the M.S. degree in electrical engineering, and the Ph.D. degree in control science from University of Minnesota, Minneapolis, USA, in 1983 and 1984, respectively. He was a post-doctoral fellow at University of Victoria, Victoria, B.C., Canada, in 1985, and a visiting assistant professor at University of Minnesota from January 1986 to April 1987. Since May 1987, he has been with the Electrical

and Computer Engineering Department, University of Victoria where he is a professor. His current research interests include analysis and design of digital filters, signal and image processing, and methods of convex optimization. He is the co-author with A. Antoniou of *Two-Dimensional Digital Filters* (Marcel Dekker, 1992) and *Practical Optimization: Algorithms and Engineering Applications* (Springer, 2007). Dr. Lu served as editor for the *Canadian Journal of Electrical and Computer Engineering* and associate editor for several journals including *IEEE Transactions on Circuits and Systems I and II*, *International Journal of Multidimensional Systems and Signal Processing*, and *Journal of Circuits, Systems, and Signal Processing*. Dr. Lu is a Life Fellow of the IEEE, a Fellow of the Engineering Institute of Canada, and a registered professional engineer in British Columbia, Canada.