

Explainable Intelligent Fault Diagnosis for Nonlinear Dynamic Systems: From Unsupervised to Supervised Learning

Hongtian Chen, *Member, IEEE*, Zhigang Liu, *Senior Member, IEEE*, Cesare Alippi, *Fellow, IEEE*,
Biao Huang, *Fellow, IEEE*, and Derong Liu, *Fellow, IEEE*

Abstract—The increased complexity and intelligence of automation systems require the development of intelligent fault diagnosis (IFD) methodologies. By relying on the concept of a suspected space, this study develops explainable data-driven IFD approaches for nonlinear dynamic systems. More in detail, we parameterize nonlinear systems through a generalized kernel representation used for system modeling and the associated fault diagnosis. An important result obtained is a unified form of kernel representations, applicable to both unsupervised and supervised learning. More importantly, through a rigorous theoretical analysis we discover the existence of a *bridge* (i.e., a bijective mapping) between some supervised and unsupervised learning-based entities. Notably, the designed IFD approaches achieve the same performance by the use of this bridge. In order to have a better understanding of the results obtained, unsupervised and supervised neural networks are chosen as the learning tools to identify generalized kernel representations and design the IFD schemes; an invertible neural network is then employed to build the bridge between them. This study is a perspective article, whose contribution lies in proposing and detailing the fundamental concepts for explainable intelligent learning methods, contributing to system modeling and data-driven IFD designs for nonlinear dynamic systems.

Index Terms—Intelligent fault diagnosis (IFD), nonlinear dynamic systems, the bridge, unsupervised learning, supervised learning, neural networks.

I. INTRODUCTION

Over the past three decades, fault diagnosis has undergone tremendous development [1]–[7]. Due to the increasing demands of safe and economic operations, it is playing an essential role in system performance evaluation and maintenance [1]. Fault diagnosis has a broad spectrum of applications, for instance, ranging from chemical processes [8], electrical systems [9], intelligent transportation [10], aerospace engineering [11] to medical imaging [12].

This work was supported by the Natural Sciences and Engineering Research Council of Canada. (*Corresponding author: Zhigang Liu and Biao Huang.*)

Hongtian Chen and Biao Huang are with the Department of Chemical and Materials Engineering, University of Alberta, Edmonton, AB T6G 1H9, Canada. Email: chtbaylor@163.com (H. Chen); biao.huang@ualberta.ca (B. Huang).

Zhigang Liu is with Institute of Rail Transit, Tongji University, Shanghai, 201804, China; School of Electrical Engineering, Southwest Jiaotong University, Chengdu, 611756, China. Email: liuzgcd@126.com.

Cesare Alippi is with the Faculty of Informatics, Università della Svizzera italiana, Lugano, 69000, Switzerland; the author is also with the Department of Electronics, Information and Bioengineering, Politecnico di Milano, Milan, 20133, Italy. Email: alippi@elet.polimi.it.

Derong Liu is with the School of Automation, Guangdong University of Technology, Guangzhou, 510006, China. Email: derong.liu@gdut.edu.cn.

Fault diagnosis techniques have undergone a dramatic evolution with the development of control theory, computer science, big data, sensor technology, machine learning, etc. [13]. They are becoming more intelligent and diversified. A unified description of these developments is so-called intelligent fault diagnosis (IFD), which can be model-based [14], signal processing-based [15], and data-driven [16].

The success of deep learning has further promoted fault diagnosis through neural networks. These IFD approaches, as mentioned in [1], can accomplish tasks similar to what a human can do. Neural networks have the ability to extract helpful features so that distinguishing faulty conditions from normal conditions becomes simpler. The accompanying problem is an additional demand of sufficient labeled data [13].

Depending on the operating range, neural networks-based IFD approaches can be divided into static and dynamic methods. In [17], a recurrent neural network was used to detect and identify faults of railway track circuits. In order to diagnose incipient interturn faults, a probabilistic neural network was adopted in [18] with consideration of the network size. By using a deep convolutional neural network, [19] proposed a diagnosis method for the loose strands of the isoelectric line. Taking the topological structure of system data into account, a graph convolutional network was developed in [20] to diagnose machine faults. Most recently, [21] proposed two fault detection schemes, where the first design is based on the finite impulse response filter using a fully connected neural network and the second one constructs a recursive residual generator using a recurrent neural network.

Despite the aforementioned achievements, much of the inherent difficulty of the IFD designs arises from:

- 1) The presence of system nonlinearity and dynamics;
- 2) Fault diagnosis cannot be simply treated as a classification problem;
- 3) Lack of interpretability in many IFD approaches, since a good representation should possess explainable posterior knowledge.

These challenges motivate the studies in this perspective article with the following four-fold contributions.

- Data-based modeling and construction of residual generators via the composite operators, enabling the learning procedures to be described quantitatively.
- A *bridge* is built for construction of the generalized kernel representations, based on which unsupervised and

supervised learning-based IFD approaches can be transferable between each other. An additional value of the bridge representation is for evaluating the performance of nonlinear IFD approaches.

- Unsupervised and supervised neural networks are employed in designing two specific IFD algorithms, whose purpose is to help us understand their fundamentals.
- In addition to the previous contribution, an invertible neural network is developed to build a bridge that allows unsupervised and supervised neural network-based IFD approach to transform between each other.

The remainder of this study is organized as follows. Section II introduces metrics, machine learning, and nonlinear dynamic systems. With the aid of composite operators, Section III is dedicated to constructing the bridge between unsupervised and supervised learning-based IFD approaches. Section IV details two specific IFD algorithms using unsupervised and supervised neural networks, followed by an invertible neural network-based bridge. Section V concludes this study and delineates research opportunities of IFD.

Notations: All notations in the article are standard. $\mathcal{M}$ denotes a metric, whose subscript signifies a specific form; $\mathcal{R}^\kappa$ represents the space of real κ -dimensional vectors; $\text{Tr}(\cdot)$ is the trace operator; $|\cdot|$ refers to the absolute value; $\circ$ is the cascade connection of multiple operators; the superscript “f” signifies the faulty conditions; $\hat{\kappa}$ is the estimate of κ ; $\text{Pr}(\cdot)$ is the probability; $\begin{pmatrix} \psi_1 & \psi_2 \\ \psi_3 & \psi_4 \end{pmatrix}$ is a composite operator (a matrix if and only if ψ_i is a linear operator) by stacking the operator ψ_i in a reasonable manner, where $i = 1, \dots, 4$.

II. PRELIMINARIES AND PROBLEM FORMULATION

A metric, also called a measure, is essential for obtaining the objective function of the machine learning approaches. Therefore, four representative metrics are introduced in this section, followed by description of the nonlinear dynamic systems of interest.

A. Metrics

Metrics are the measures that can quantitatively assess, compare, and track performance [22]. In the following, several metrics are introduced for the purposes like constructing residual signals, defining test statistics, and quantifying the influences of the unknown faults.

Consider two variables $\psi_1 \in \mathcal{R}^{k_{\psi_1}}$ and $\psi_2 \in \mathcal{R}^{k_{\psi_2}}$. If a metric $\mathcal{M}$ is chosen as the Euclidean distance, one has

$$\mathcal{M}_{\text{euc}} = \|\psi_1 - \psi_2\|_2, \quad (1)$$

where $k_{\psi_1} = k_{\psi_2}$. By introducing a covariance matrix Σ_ψ , the Mahalanobis distance can be defined by

$$\mathcal{M}_{\text{mah}}^2 = (\psi_1 - \psi_2)^T \Sigma_\psi^{-1} (\psi_1 - \psi_2), \quad (2)$$

in which ψ_1 and ψ_2 are from the same distribution with covariance Σ_ψ . By setting $\Sigma_\psi^{-1} = \mathbf{I}$, (3) becomes

$$\mathcal{M}_{\text{spe}}^2 = (\psi_1 - \psi_2)^T (\psi_1 - \psi_2), \quad (3)$$

which is called squared prediction error (SPE). A metric can also be defined by the correlation such as [23]

$$\mathcal{M}_{\text{cor}}(\psi_1, \psi_2) = k_{\psi_1} - \text{Tr} \left(|\Sigma_{\psi_1}^{-1/2} \Sigma_{\psi_1, \psi_2} \Sigma_{\psi_2}^{-1/2}| \right), \quad (4a)$$

$$\mathcal{M}_{\text{cor}}(\psi_2, \psi_1) = k_{\psi_2} - \text{Tr} \left(|\Sigma_{\psi_2}^{-1/2} \Sigma_{\psi_2, \psi_1} \Sigma_{\psi_1}^{-1/2}| \right), \quad (4b)$$

where Σ_{ψ_1, ψ_2} is the covariance matrix between ψ_1 and ψ_2 , and k_{ψ_1} may not equal to k_{ψ_2} .

B. Unsupervised and Supervised Machine Learning

Depending upon linear or nonlinear mappings, machine learning is generally divided into linear and nonlinear approaches. An important step is to evaluate performance through the defined matrices (such as $\mathcal{M}_{\text{euc}}$, $\mathcal{M}_{\text{mah}}$, $\mathcal{M}_{\text{spe}}$, and $\mathcal{M}_{\text{cor}}$), and the subsequent step will be devoted to revealing the relationship between unsupervised and supervised machine learning approaches.

Define $\psi_2 = \mathcal{S}\psi_3$, where $\mathcal{S}$ represents the mapping of supervised learning methods. Considering the objective function to minimize (1), the following relation holds

$$\begin{aligned} \min \mathcal{M}_{\text{euc}} &= \min \left\| \begin{bmatrix} \mathbf{I} & -\mathcal{S} \end{bmatrix} \begin{bmatrix} \psi_1 \\ \psi_3 \end{bmatrix} \right\|_2 \\ &= \min \left\| \underbrace{\mathcal{P} \circ \begin{bmatrix} \mathbf{I} & -\mathcal{S} \end{bmatrix}}_{\text{unsupervised learning}} \begin{bmatrix} \psi_1 \\ \psi_3 \end{bmatrix} \right\|_2, \end{aligned} \quad (5)$$

where $\mathcal{P}$ is any (linear or nonlinear) operator that does not cause information loss or change the minimization of (5). In (1) and (5), ψ_1 is a reference signal; ψ_2 and ψ_3 are the output and input of supervised learning models, respectively. By introducing $\mathcal{P}$, $\mathcal{P} \circ [\mathbf{I} \quad -\mathcal{S}]$ is a composite operator whose input is $[\psi_1^T \quad \psi_3^T]^T$. Instead of optimizing $\mathcal{S}$, we can treat $\mathcal{P} \circ [\mathbf{I} \quad -\mathcal{S}]$ as a whole item to learn. In this manner, the objective function given in (5) can be minimized through an unsupervised learning approach. Similarly, the objective functions defined by other metrics (including $\mathcal{M}_{\text{mah}}$, $\mathcal{M}_{\text{spe}}$, and $\mathcal{M}_{\text{cor}}$) can also be achieved.

It is seen that, through simple mathematical derivations, an objective function can be achieved by using both the supervised and unsupervised approaches.

Remark 1. *The concept behind the aforementioned transformations is a bridge linking the unsupervised to supervised machine learning methods. This fact motivates the current study to focus on the development of a unified framework of the IFD and related parameter-identification approaches for nonlinear systems.* ∇

C. Nonlinear Dynamic Systems and Objectives

Consider a nonlinear dynamic system driven by

$$\begin{aligned} \mathbf{x}(k+1) &= \phi(\mathbf{x}(k), \mathbf{u}(k), \mathbf{w}(k)), \\ \mathbf{y}(k) &= \mathbf{v}(\mathbf{x}(k), \mathbf{u}(k)) + \mathbf{v}(k), \end{aligned} \quad (6)$$

where k is the discrete time index; $\mathbf{x}(k) \in \mathcal{R}^{k_x}$, $\mathbf{u}(k) \in \mathcal{R}^{k_u}$, and $\mathbf{y}(k) \in \mathcal{R}^{k_y}$ are the state, input, and output vectors,

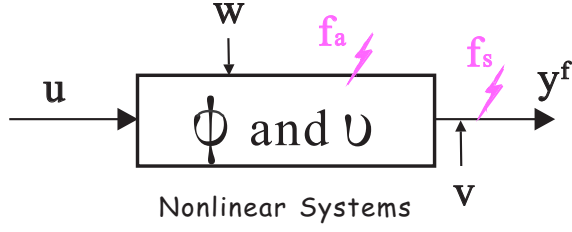


Fig. 1: Nonlinear dynamic systems with actuator and sensor faults.

respectively; $\mathbf{w}(k)$ and $\mathbf{v}(k)$ are random noise variables; $\phi(\cdot)$ and $\psi(\cdot)$ are continuous nonlinear mappings.

Both the actuator and sensor faults affect (6). Mathematically, the nonlinear system with faults becomes

$$\begin{aligned} \mathbf{x}^f(k+1) &= \phi(\mathbf{x}(k), \mathbf{u}(k), \mathbf{w}(k), \mathbf{f}_a(k)), \\ \mathbf{y}^f(k) &= \psi(\mathbf{x}(k), \mathbf{u}(k), \mathbf{f}_a(k)) + \mathbf{f}_s(k) + \mathbf{v}(k), \end{aligned} \quad (7)$$

where $\mathbf{f}_a$ and $\mathbf{f}_s$ represent actuator and sensor faults, respectively. The schematic of (7) is given in Fig. 1. Now we introduce a lemma (given in [24]) that constitutes another foundation of this study.

Lemma 1. *An operator $\mathcal{K}$ is called the generalized kernel representation for the nonlinear system (6) if for $\mathbf{w} = \mathbf{0}$ and $\mathbf{v} = \mathbf{0}$, the following relationship holds*

$$\mathcal{K}(z) \begin{bmatrix} \mathbf{u}(z) \\ \mathbf{y}(z) \end{bmatrix} = \mathbf{0}, \quad (8)$$

for an initial system state $\mathbf{x}(0)$ and a given $\mathbf{u}$.

In (8), z signifies the discrete variables in the z domain. Motivated by [16], [21], [23], and [25], this study focuses on IFD and its related parameter identification for nonlinear dynamic systems. The three objectives are casted as follows.

- To construct a unified IFD framework that can include both unsupervised and supervised learning-based approaches;
- To design two IFD approaches using the unsupervised and supervised neural networks, respectively;
- To quantify the fault influences in each situation, whose purpose is to mine the interpretability of IFD methods.

III. A UNIFIED IFD FRAMEWORK: FROM SUPERVISED TO UNSUPERVISED LEARNING

By using the composite operators and suspected spaces (see [23]), this section will present data-driven implementation of residual generators using both unsupervised and supervised learning, with a focus on construction of the bridge between unsupervised and supervised learning-based approaches.

A. Data-based Modeling

Because both $\mathbf{x}(k)$ and the innovation from $\mathbf{x}(k-1)$ to $\mathbf{x}(k)$ are unknown, the extended form of (6) is usually adopted for system identification and fault diagnosis. For obtaining

the extended form, several notations, referring to data and operators, are introduced:

$$\begin{aligned} \phi_{\wp_1 \wp_2 \wp_2}(k) &: \wp_2(k) \in \mathcal{R}^{k_{\wp_2}} \rightarrow \wp_1(k+1) \in \mathcal{R}^{k_{\wp_1}}, \\ \psi_{\wp_1 \wp_2 \wp_2}(k) &: \wp_2(k) \in \mathcal{R}^{k_{\wp_2}} \rightarrow \wp_1(k) \in \mathcal{R}^{k_{\wp_1}}, \\ \phi_{\wp_1 \wp_1}^k &= \phi_{\wp_1 \wp_1}^{k-1} \circ \phi_{\wp_1 \wp_1} = \phi_{\wp_1 \wp_1}^{k-2} \circ \phi_{\wp_1 \wp_1} \circ \phi_{\wp_1 \wp_1}, \end{aligned} \quad (9)$$

and

$$\wp_s(k) = [\wp^T(k) \cdots \wp^T(k+s)]^T \in \mathcal{R}^{(s+1)k_{\wp}}. \quad (10)$$

In (9), $\wp_1$ and $\wp_2$ are vectors; $\phi_{\wp_1 \wp_2}$ and $\psi_{\wp_1 \wp_2}$ are linear or nonlinear operators, respectively; $\phi_{\wp_1 \wp_1}^k$ represents a high-order composite operator. In (10), s is the stack length. Furthermore, $\wp_1$, $\wp_2$, and $\wp$ can be replaced by any variable in (6) and (7); $\phi_{\cdot}$ and $\psi_{\cdot}$ can also be replaced by any operator.

Remark 2. *In order to parameterize nonlinear dynamic systems, the composite operators defined in (9), similar to Koopman operators [26], simplify the treatment processes for obtaining an equivalent system representation [27].* ∇

By using the notation given in (9), we rewrite (6) as

$$\mathbf{x}(k+1) = \underbrace{(\phi_{\mathbf{x}\mathbf{x}} \quad \phi_{\mathbf{x}\mathbf{u}} \quad \phi_{\mathbf{x}\mathbf{w}})}_{\phi} \begin{bmatrix} \mathbf{x}(k) \\ \mathbf{u}(k) \\ \mathbf{w}(k) \end{bmatrix}, \quad (11a)$$

$$\mathbf{y}(k) = \underbrace{(\psi_{\mathbf{y}\mathbf{x}} \quad \psi_{\mathbf{y}\mathbf{u}} \quad \psi_{\mathbf{y}\mathbf{v}})}_{\psi} \begin{bmatrix} \mathbf{x}(k) \\ \mathbf{u}(k) \\ \mathbf{v}(k) \end{bmatrix}, \quad (11b)$$

where $\psi_{\mathbf{y}\mathbf{v}} = \mathbf{I}$. Similar to the parity space in [2] and [11], the data-based system model can be described as

$$\mathbf{y}_s(k) = (\Upsilon_{\mathbf{x}} \quad \Upsilon_{\mathbf{u}}) \begin{bmatrix} \mathbf{x}(k) \\ \mathbf{u}_s(k) \end{bmatrix} + \Upsilon_{\mathbf{w}} \mathbf{w}_s(k) + \mathbf{v}_s(k), \quad (12)$$

in which the composite operators $\Upsilon_{\mathbf{x}}$, $\Upsilon_{\mathbf{u}}$, and $\Upsilon_{\mathbf{w}}$ are

$$\Upsilon_{\mathbf{x}} := \begin{pmatrix} \psi_{\mathbf{y}\mathbf{x}} \\ \vdots \\ \psi_{\mathbf{y}\mathbf{x}} \phi_{\mathbf{x}\mathbf{x}}^s \end{pmatrix} : \mathcal{R}^{k_x} \rightarrow \mathcal{R}^{(s+1)k_y}, \quad (13)$$

$$\Upsilon_{\mathbf{u}} := \begin{pmatrix} \psi_{\mathbf{y}\mathbf{u}} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ \psi_{\mathbf{y}\mathbf{x}} \phi_{\mathbf{x}\mathbf{x}}^{s-1} \phi_{\mathbf{x}\mathbf{u}} & \cdots & \psi_{\mathbf{y}\mathbf{u}} \end{pmatrix} : \mathcal{R}^{(s+1)k_u} \rightarrow \mathcal{R}^{(s+1)k_y}, \quad (14)$$

$$\Upsilon_{\mathbf{w}} := \begin{pmatrix} \mathbf{0} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ \psi_{\mathbf{y}\mathbf{x}} \phi_{\mathbf{x}\mathbf{x}}^{s-1} \phi_{\mathbf{x}\mathbf{w}} & \cdots & \mathbf{0} \end{pmatrix} : \mathcal{R}^{(s+1)k_w} \rightarrow \mathcal{R}^{(s+1)k_y}. \quad (15)$$

In the data-based model (12), the state $\mathbf{x}(k)$ is generally unknown and needs to be estimated by using the past data. For example, the studies [21], [28], [29] suggest to estimate $\mathbf{y}$ through data in the past moving window, i.e.,

$$\mathbf{e}(k) = \mathbf{y}(k) - \hat{\mathbf{y}}(k), \hat{\mathbf{y}}(k) = \mathbf{v}(\hat{\mathbf{x}}(k), \mathbf{u}(k)), \quad (16a)$$

$$\hat{\mathbf{x}}(k+1) = \phi(\hat{\mathbf{x}}(k), \mathbf{u}(k)) + \ell(\mathbf{y}(k) - \hat{\mathbf{y}}(k)), \quad (16b)$$

in which $\ell : \mathcal{R}^{k_y} \rightarrow \mathcal{R}^{k_x}$ signifies the (unknown but existing) projection from $\mathbf{e}$ to $\hat{\mathbf{x}}$. Therefore, one can obtain that

$$\begin{aligned} \hat{\mathbf{x}}(k) &= \phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}}^{s_p+1} \hat{\mathbf{x}}(k-s_p-1) + [\mathcal{L}_{\mathbf{p},\mathbf{u}} \quad \mathcal{L}_{\mathbf{p},\mathbf{y}}] \mathbf{z}_{\mathbf{p}}(k), \\ \phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}}^{s_p+1} \hat{\mathbf{x}}(k-s_p-1) &\approx \mathbf{0}, \\ \Rightarrow \hat{\mathbf{x}}(k) &\approx \underbrace{[\mathcal{L}_{\mathbf{p},\mathbf{u}} \quad \mathcal{L}_{\mathbf{p},\mathbf{y}}]}_{\mathcal{L}_{\mathbf{p}}} \mathbf{z}_{\mathbf{p}}(k), \end{aligned} \quad (17)$$

where $\mathcal{L}_{\mathbf{p},\mathbf{u}}$ and $\mathcal{L}_{\mathbf{p},\mathbf{y}}$ are the composite operators similar to (14)-(15); the past stacked vector $\mathbf{z}_{\mathbf{p}}(k)$ is defined by

$$\begin{aligned} \mathbf{z}_{\mathbf{p}}(k) &= \begin{bmatrix} \mathbf{u}_{\mathbf{p}}(k) \\ \mathbf{y}_{\mathbf{p}}(k) \end{bmatrix}, \mathbf{u}_{\mathbf{p}}(k) = \mathbf{u}_{s_p}(k-s_p-1), \\ \mathbf{y}_{\mathbf{p}}(k) &= \mathbf{y}_{s_p}(k-s_p-1). \end{aligned} \quad (18)$$

Combining (12) with (17) yields an equivalent model:

$$\begin{aligned} \mathbf{y}_s(k) &= \Upsilon_{\mathbf{x}} \mathcal{L}_{\mathbf{p}} \mathbf{z}_{\mathbf{p}}(k) + \Upsilon_{\mathbf{u}} \mathbf{u}_s(k) + \Upsilon_{\mathbf{e}} \mathbf{e}_s(k) \\ &= (\Upsilon_{\mathbf{x}} \mathcal{L}_{\mathbf{p}} \quad \Upsilon_{\mathbf{u}}) \begin{bmatrix} \mathbf{z}_{\mathbf{p}}(k) \\ \mathbf{u}_s(k) \end{bmatrix} + \Upsilon_{\mathbf{e}} \mathbf{e}_s(k) \end{aligned} \quad (19)$$

where $\mathbf{e}_s$ contains the influences caused by $\mathbf{w}$ and $\mathbf{v}$, and $\Upsilon_{\mathbf{e}}$ has the following form:

$$\Upsilon_{\mathbf{e}} = \begin{pmatrix} \mathbf{I} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}}^{s-1} \ell_{\hat{\mathbf{x}}\mathbf{e}} & \cdots & \mathbf{I} \end{pmatrix}. \quad (20)$$

Remark 3. By the use of the composite operator, (19) details a generalized nonlinear predictor from $\mathbf{z}_{\mathbf{p}}$ and $\mathbf{u}_s$ to the current system output $\mathbf{y}_s$, while resembling a linear model. It makes system modeling possible and explainable, especially achieving a consensus among different IFD approaches using both unsupervised and supervised learning. ∇

B. Data-driven Implementation of Residual Generators

On the basis of Lemma 1, the following corollary is obtained, whose embryonic development credits to [21].

Corollary 1. An operator $\mathcal{K}_{s_p+s}$ is called the data-driven generalized kernel representation for (6) if for $\mathbf{w} = \mathbf{0}$ and

$\mathbf{v} = \mathbf{0}$, the following relationship holds

$$\underbrace{(\mathcal{K}_{\mathbf{z}}^{\text{unsp}} \quad \mathcal{K}_{\mathbf{u}}^{\text{unsp}} \quad \mathcal{K}_{\mathbf{y}}^{\text{unsp}})}_{\mathcal{K}_{s_p+s}^{\text{unsp}}} \begin{bmatrix} \mathbf{z}_{\mathbf{p}} \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} = \mathbf{0}, \quad (21a)$$

$$\text{or } \underbrace{(\mathcal{K}_{s_p+s}^{\text{sp}} - \mathbf{I})}_{\mathcal{K}_{s_p+s}^{\text{sp}}} \begin{bmatrix} \mathbf{z}_{\mathbf{p}} \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} = \mathbf{0}, \quad (21b)$$

for an initial system state $\mathbf{x}(0)$ and an arbitrary $\mathbf{u}_s$, where the superscripts “unsp” and “sp” signify the unsupervised and supervised learning scenarios, respectively.

The obtained $\mathcal{K}_{s_p+s}$, by using the unsupervised or supervised learning methods, can be directly applied in constructing a unified residual generator such that

$$\mathbf{r}_s(k) = \mathcal{K}_{s_p+s} \begin{bmatrix} \mathbf{z}_{\mathbf{p}}(k) \\ \mathbf{u}_s(k) \\ \mathbf{y}_s(k) \end{bmatrix} \in \mathcal{R}^{(s+1)k_y}, \quad (22)$$

in fault-free conditions; $\mathcal{K}_{s_p+s}$ can be replaced by $\mathcal{K}_{s_p+s}^{\text{unsp}}$ and $\mathcal{K}_{s_p+s}^{\text{sp}}$ provided in Corollary 1. When (6) is affected by faults, $\mathbf{r}_s(k)$ given in (22) becomes

$$\mathbf{r}_s^f(k) = \mathbf{r}_s(k) + \mathbf{f}_{\mathbf{a},\text{term}} + \mathbf{f}_{\mathbf{s},\text{term}}, \quad (23)$$

where $\mathbf{f}_{\mathbf{a},\text{term}}$ and $\mathbf{f}_{\mathbf{s},\text{term}}$ are the actuator and sensor fault-related terms, respectively; their forms can be obtained via composite operators. Then the test statistics [1], such as T^2 , can be defined on the residual signal $\mathbf{r}_s$, i.e.,

$$T^2(\mathbf{r}_s(k)) = \mathbf{r}_s^T(k) \Sigma_{\mathbf{r}_s}^{-1} \mathbf{r}_s(k), \quad \Sigma_{\mathbf{r}_s}^{-1} = \mathbb{E}(\mathbf{r}_s(k) \mathbf{r}_s^T(k)). \quad (24)$$

Correspondingly, the threshold is determined by the chosen confidence level α :

$$J_{th} \leftarrow \Pr(T^2(\mathbf{r}_s(k)) > J_{th}) = \alpha. \quad (25)$$

The following theorem presents the essence of this study, sketching *the bridge* (i.e., *one-to-one mapping*) between unsupervised and supervised learning-based parameter identification, together with the IFD approaches, for nonlinear systems.

Theorem 1. Consider the two residual generators defined by

$$\mathbf{r}_s^{\text{unsp}}(k) = \mathcal{K}_{s_p+s}^{\text{unsp}} \begin{bmatrix} \mathbf{z}_{\mathbf{p}}(k) \\ \mathbf{u}_s(k) \\ \mathbf{y}_s(k) \end{bmatrix}, \quad \mathbf{r}_s^{\text{sp}}(k) = \mathcal{K}_{s_p+s}^{\text{sp}} \begin{bmatrix} \mathbf{z}_{\mathbf{p}}(k) \\ \mathbf{u}_s(k) \\ \mathbf{y}_s(k) \end{bmatrix} \quad (26)$$

where $\mathcal{K}_{s_p+s}^{\text{unsp}}$ and $\mathcal{K}_{s_p+s}^{\text{sp}}$ are defined by

$$\mathcal{K}_{s_p+s}^{\text{unsp}} = [\mathbf{0} \quad \mathbf{0} \quad \mathbf{I}] - \mathcal{U}_{\mathbf{y}}, \quad (27a)$$

$$\mathcal{K}_{s_p+s}^{\text{sp}} = (-\Upsilon_{\mathbf{x}} \mathcal{L}_{\mathbf{p}} \quad -\Upsilon_{\mathbf{u}} \quad \mathbf{I}), \quad (27b)$$

in which $\mathcal{U}_y$ is a segment of unsupervised learning $\mathcal{U}$:

$$\begin{bmatrix} \hat{\mathbf{z}}_p(k) \\ \hat{\mathbf{u}}_s(k) \\ \hat{\mathbf{y}}_s(k) \end{bmatrix} = \underbrace{(\mathcal{U}_z \quad \mathcal{U}_u \quad \mathcal{U}_y)}_{\mathcal{U}} \begin{bmatrix} \mathbf{z}_p(k) \\ \mathbf{u}_s(k) \\ \mathbf{y}_s(k) \end{bmatrix}. \quad (28)$$

There exists a nonlinear operator $\mathcal{P}_{\text{sp/uns}}^{\text{sp}}$ such that

$$\mathbf{r}_s^{\text{sp}}(k) = \mathcal{P}_{\text{sp/uns}}^{\text{sp}} \mathbf{r}_s^{\text{uns}}(k), \quad (29)$$

and the inverse function of $\mathcal{P}_{\text{sp/uns}}^{\text{sp}}$, denoted as $\mathcal{P}_{\text{uns/sp}}^{\text{sp}} = \mathcal{P}_{\text{sp/uns}}^{\text{sp}-1}$, must exist such that

$$\mathbf{r}_s^{\text{uns}}(k) = \mathcal{P}_{\text{uns/sp}}^{\text{sp}} \mathbf{r}_s^{\text{sp}}(k). \quad (30)$$

Proof: The complete proof is given in Appendix A. ■

The following remark is made to set forth contributions of Theorem 1.

Remark 4. With the help of the bridge, Theorem 1 provides us with an elegant way to evaluate the performance of both unsupervised and supervised learning-based IFD approaches in the sense of fault-detection power. ▽

Based on the metric $\mathcal{M}_{\text{euc}}^2$, the suspected space is used to design IFD approaches, whose purpose is to gain a more in-depth understanding of Theorem 1.

C. IFD Using Unsupervised Learning

In order to develop the unsupervised learning-based IFD approaches, a suspected space, denoted as $\mathcal{Y}^{\text{uns}}^{\text{sp}}$, is defined according to $\mathcal{U}_y$ in the following definition.

Definition 1. Given any $\mathcal{U}$ in (28), $\mathcal{Y}_s^{\text{uns}}^{\text{sp}}$ obtained by

$$\mathcal{Y}_s^{\text{uns}}^{\text{sp}} = \mathcal{U}_y [\mathbf{z}_p^T \quad \mathbf{u}_s^T \quad \mathbf{y}_s^T]^T, \quad (31)$$

spans a space $\mathcal{Y}^{\text{uns}}^{\text{sp}}$ that is called the suspected space of (6).

Corresponding to $\mathcal{Y}^{\text{uns}}^{\text{sp}}$, $\mathcal{Y}$ is called the measurement space, where $\mathbf{y}_s \in \mathcal{Y} \subset \mathcal{R}^{(s+1)k_y}$. Based on the metric defined in (1), the objective function of IFD approaches using unsupervised learning can be formulated as

$$\min \mathcal{M}_{\text{euc}}(\mathbf{y}_s, \mathcal{Y}_s^{\text{uns}}^{\text{sp}}), \quad (32a)$$

$$\Rightarrow \mathcal{U}_y^* := \arg \min_{\mathcal{U}_y} \mathcal{M}_{\text{euc}}(\mathbf{y}_s, \mathcal{Y}_s^{\text{uns}}^{\text{sp}}), \quad (32b)$$

where the superscript “*” signifies the best solution. It is of interest to find that the residual space, denoted as $\mathcal{E}_s^{\text{uns}}^{\text{sp}}$, can be obtained based on $\mathcal{U}_y^*$:

$$\begin{aligned} \mathbf{e}_s &= \mathbf{y}_s - \mathcal{Y}_s^{\text{uns},*}, \mathbf{e}_s \in \mathcal{E}_s^{\text{uns}}^{\text{sp}} \subset \mathcal{R}^{(s+1)k_y}, \\ \mathcal{Y}_s^{\text{uns},*} &= \mathcal{U}_y^* [\mathbf{z}_p^T \quad \mathbf{u}_s^T \quad \mathbf{y}_s^T]^T \in \mathcal{Y}^{\text{uns},*}, \end{aligned} \quad (33)$$

which satisfies

$$\mathcal{E}_s^{\text{uns}}^{\text{sp}} \perp \mathcal{Y}^{\text{uns},*}, \quad \mathcal{E}_s^{\text{uns}}^{\text{sp}} (\approx) \perp \mathcal{Y}. \quad (34)$$

Theorem 2. Considering a nonlinear system (6), its generalized kernel representation is defined by $\mathcal{K}_{s_p+s}^{\text{uns}} = [\mathbf{0} \quad \mathbf{0} \quad \mathbf{I}] - \mathcal{U}_y^*$, where $\mathcal{U}_y^*$ is obtained through (32b). For the faults $\mathbf{f}_a(k)$

and $\mathbf{f}_s(k)$ occurring from the k -th time instant, $\mathbf{r}_s^{\text{uns}}(k)$ given in (26) becomes

$$\begin{aligned} \mathbf{r}_s^{\text{uns},f}(k) &= \mathbf{e}_s(k) + \mathbf{f}_{a,\text{term}}^{\text{uns}}(k) + \mathbf{f}_{s,\text{term}}^{\text{uns}}(k), \\ \mathbf{f}_{a,\text{term}}^{\text{uns}}(k) &= \underline{\Upsilon}_{f_a} \mathbf{f}_{a,s}(k), \quad \mathbf{f}_{s,\text{term}}^{\text{uns}}(k) = \underline{\Upsilon}_{f_s} \mathbf{f}_{s,s}, \end{aligned} \quad (35)$$

where $\mathbf{f}_{a,\text{term}}^{\text{uns}}$ and $\mathbf{f}_{s,\text{term}}^{\text{uns}}$ have the following forms:

$$\underline{\Upsilon}_{f_a} = \underline{\Upsilon}_u, \quad \underline{\Upsilon}_{f_s} = \begin{pmatrix} \mathbf{I} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ -\mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{A}_{\ell}^{s-1} \ell_{\hat{\mathbf{x}}\mathbf{e}} & \cdots & \mathbf{I} \end{pmatrix}. \quad (36)$$

Then T^2 defined on $\mathbf{r}_s^{\text{uns},f}(k)$ according to

$$T^2(\mathbf{r}_s^{\text{uns},f}(k)) = \mathbf{r}_s^{\text{uns},f,T}(k) \Sigma_{r_s}^{-1} \mathbf{r}_s^{\text{uns},f}(k) \quad (37)$$

has the optimal fault-detection power.

Proof: When (6) is fault-free, we have

$$\mathbf{r}_s^{\text{uns}}(k) = \mathbf{e}_s(k), \quad \Sigma_{r_s}^{\text{uns}} = \mathbb{E}(\mathbf{e}_s(k) \mathbf{e}_s^T(k)). \quad (38)$$

Since $\mathbf{f}_a$ and $\mathbf{f}_s$ respectively represent the actuator and sensor faults, it is reasonable to adopt

$$\phi_{\mathbf{x}f_a} = \phi_{\mathbf{x}u}, \quad \mathbf{v}_{\mathbf{y}f_a} = \mathbf{v}_{\mathbf{y}u}, \quad \text{and} \quad \mathbf{v}_{\mathbf{y}f_s} = \mathbf{I}. \quad (39)$$

When (6) is affected by $\mathbf{f}_a$ and $\mathbf{f}_s$, $\hat{\mathbf{x}}^f(k)$ and $\mathbf{y}^f(k)$ in (16a)-(16b) can be expressed by

$$\begin{aligned} \hat{\mathbf{x}}^f(k+1) &= \phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}} \hat{\mathbf{x}}(k) + \phi_{\hat{\mathbf{x}}u} \mathbf{u}(k) + \phi_{\hat{\mathbf{x}}u} \mathbf{f}_a(k) \\ &\quad + \ell(\mathbf{y}^f(k) - \hat{\mathbf{y}}(k) - \mathbf{f}_s), \\ \mathbf{y}^f(k) &= \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \hat{\mathbf{x}}(k) + \mathbf{v}_{\hat{\mathbf{y}}u} \mathbf{u}(k) + \mathbf{v}_{\hat{\mathbf{y}}u} \mathbf{f}_a(k) \\ &\quad + \mathbf{v}_{\mathbf{y}f_s} \mathbf{f}_s(k) + \mathbf{e}(k). \end{aligned} \quad (40)$$

Combining it with (A.2) yields (35) with $\underline{\Upsilon}_{f_a}$ and $\underline{\Upsilon}_{f_s}$ described by (36). Furthermore, taking expectation of (37) obtains

$$\begin{aligned} \mathbb{E}(T^2(\mathbf{r}_s^{\text{uns},f}(k))) &= T^2(\mathbf{e}_s) + \mathbf{f}_{a,\text{term}}^T \Sigma_{r_s}^{-1} \mathbf{f}_{a,\text{term}} \\ &\quad + \mathbf{f}_{f,\text{term}}^T \Sigma_{r_s}^{-1} \mathbf{f}_{f,\text{term}}, \end{aligned} \quad (41)$$

because in general $\mathbf{e}_s(k)$, $\mathbf{f}_{a,\text{term}}(k)$, and $\mathbf{f}_{f,\text{term}}(k)$ are independent. As proved in [23], $\mathcal{U}_y^*$ minimizes $\Sigma_{r_s}^{\text{uns}}$ via (32b); thus the last two terms in (41) are maximized to the greatest extent possible. It means $\mathcal{U}_y^*$, a segment of $\mathcal{U}^*$, can achieve optimal performance in detecting $\mathbf{f}_a$ and $\mathbf{f}_s$.

The theorem is proved. ■

D. IFD Using Supervised Learning

In parallel to $\mathcal{Y}^{\text{uns}}^{\text{sp}}$, another suspected space $\mathcal{Y}^{\text{sp}}$ using supervised learning can be defined according to $\mathcal{S}$ as follows.

Definition 2. Given any $\mathcal{S}$ in (A.8), $\mathcal{Y}_s^{\text{sp}}$ obtained by

$$\mathcal{Y}_s^{\text{sp}} = \mathcal{S} [\mathbf{z}_p^T \quad \mathbf{u}_s^T]^T, \quad (42)$$

spans a space $\mathcal{Y}^{\text{sp}}$ that is also the suspected space of (6), where $\mathcal{S}$ is defined by

$$\mathcal{S} = (\mathcal{S}_z \quad \mathcal{S}_u) := (\Upsilon_x \mathcal{L}_p \quad \Upsilon_u). \quad (43)$$

Similar to (32a)-(32b), the optimal $\mathcal{S}^*$ can be obtained via

$$\min \mathcal{M}_{\text{euc}}(\mathbf{y}_s, \mathbf{y}_s^{\text{sp}}), \quad (44a)$$

$$\Rightarrow \mathcal{S}^* := \arg \min_{\mathcal{S}} \mathcal{M}_{\text{euc}}(\mathbf{y}_s, \mathbf{y}_s^{\text{sp}}), \quad (44b)$$

based on which $\mathbf{y}^{\text{sp},*}$ can be obtained.

By the use of $\mathcal{S}^*$, the following theorem is derived for IFD using supervised learning.

Theorem 3. *Considering a nonlinear system (6), its generalized kernel representation is defined by $\mathcal{K}_{s_p+s}^{\text{sp}} = (\mathcal{S}^* - \mathbf{I})$, where $\mathcal{S}^*$ is obtained through (44b). For the faults $\mathbf{f}_a(k)$ and $\mathbf{f}_s(k)$ occurring from the k -th time instant, $\mathbf{r}_s^{\text{sp}}(k)$ given in (26) becomes*

$$\begin{aligned} \mathbf{r}_s^{\text{sp},f}(k) &= \Upsilon_e \mathbf{e}_s(k) + \mathbf{f}_{a,\text{term}}^{\text{sp}}(k) + \mathbf{f}_{s,\text{term}}^{\text{sp}}(k), \\ \mathbf{f}_{a,\text{term}}^{\text{sp}}(k) &= \Upsilon_{f_a} \mathbf{f}_{a,s}(k), \quad \mathbf{f}_{s,\text{term}}^{\text{sp}}(k) = \Upsilon_{f_s} \mathbf{f}_{s,s}, \end{aligned} \quad (45)$$

where $\mathbf{f}_{a,\text{term}}^{\text{sp}}$ and $\mathbf{f}_{s,\text{term}}^{\text{sp}}$ have the following forms:

$$\Upsilon_{f_a} = \Upsilon_u, \quad \Upsilon_{f_s} = \mathbf{I}. \quad (46)$$

Then T^2 defined on $\mathbf{r}_s^{\text{sp},f}(k)$ according to

$$T^2(\mathbf{r}_s^{\text{sp},f}(k)) = \mathbf{r}_s^{\text{sp},f,T}(k) \Sigma_{\mathbf{r}_s}^{-1} \mathbf{r}_s^{\text{sp},f}(k) \quad (47)$$

has the optimal fault-detection power.

Proof: The proof is similar to Theorem 2. Owing to space constraints, the detail is omitted here.

In addition, another sketch of the proof for the theorem is presented as follows. By using the bridge $\mathcal{P}_{\text{sp}/\text{unsp}}$ given in Theorem 1, simple mathematical manipulations can yield

$$T^2(\mathbf{r}_s^{\text{sp}}(k)) = T^2(\mathbf{r}_s^{\text{unsp}}(k)), \quad (48)$$

and (46). Furthermore, it can also be verified

$$T^2(\mathbf{r}_s^{\text{sp},f}(k)) = T^2(\mathbf{r}_s^{\text{unsp},f}(k)), \quad (49)$$

which has the same (i.e., optimal) performance of fault detection as $T^2(\mathbf{r}_s^{\text{unsp},f}(k))$, and thus completes this proof. ■

Remark 5. *As presented in Theorems 2 and 3, performance evaluation of the two proposed IFD frameworks for nonlinear systems becomes possible. The main reason is that the difference between $\mathbf{y}_s(k)$ and its estimation in $\mathbf{y}^{\text{unsp},*}$ and $\mathbf{y}^{\text{sp},*}$ can be measured through quantitative metrics.* ▽

E. Notes on the Bridge

In what follows, several remarks (including fault features and geometric interpretation) and perspectives (i.e., a more general version) are made to set forth contributions and essences of the bridge provided in Theorem 1.

1) *Fault features:* Along with Theorems 2 and 3, the fault features, corresponding to $\mathbf{y}^{\text{unsp},*}$ and $\mathbf{y}^{\text{sp},*}$, are

$$\mathbf{r}_s^{\text{unsp},f} = \mathbf{e}_s + \mathbf{f}_{a,\text{term}}^{\text{unsp}} + \mathbf{f}_{s,\text{term}}^{\text{unsp}}, \quad (50a)$$

$$\mathbf{r}_s^{\text{sp},f} = \mathcal{P}_{\text{sp}/\text{unsp}} \mathbf{e}_s + \mathbf{f}_{a,\text{term}}^{\text{sp}} + \mathbf{f}_{s,\text{term}}^{\text{sp}}. \quad (50b)$$

It is interesting to verify

$$\begin{aligned} \mathcal{P}_{\text{sp}/\text{unsp}} \mathbf{f}_{a,\text{term}}^{\text{unsp}} &= \mathbf{f}_{a,\text{term}}^{\text{sp}}, \\ \mathcal{P}_{\text{unsp}/\text{sp}} \mathbf{f}_{a,\text{term}}^{\text{sp}} &= \mathcal{P}_{\text{sp}/\text{unsp}}^{-1} \mathbf{f}_{a,\text{term}}^{\text{sp}} = \mathbf{f}_{a,\text{term}}^{\text{unsp}}, \end{aligned} \quad (51)$$

for $\mathbf{f}_{a,\text{term}}$, and

$$\begin{aligned} \mathcal{P}_{\text{sp}/\text{unsp}} \mathbf{f}_{s,\text{term}}^{\text{unsp}} &= \mathbf{f}_{s,\text{term}}^{\text{sp}}, \\ \mathcal{P}_{\text{unsp}/\text{sp}} \mathbf{f}_{s,\text{term}}^{\text{sp}} &= \mathbf{f}_{s,\text{term}}^{\text{unsp}}, \end{aligned} \quad (52)$$

for $\mathbf{f}_{s,\text{term}}$. Combining (50a)-(50b) with (51) and (52) yields

$$\mathcal{P}_{\text{sp}/\text{unsp}} \mathbf{r}_s^{\text{unsp},f} = \mathbf{r}_s^{\text{sp},f} = \mathcal{P}_{\text{sp}/\text{unsp}} \circ \mathcal{P}_{\text{sp}/\text{unsp}}^{-1} \mathbf{r}_s^{\text{sp},f}, \quad (53)$$

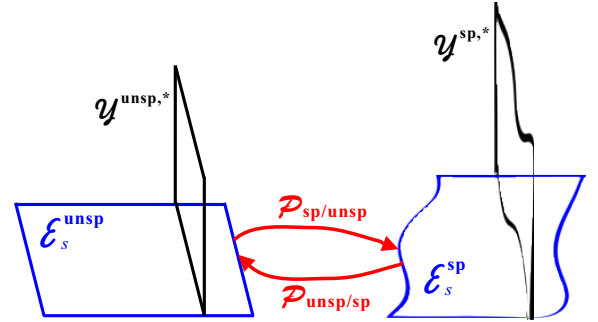
which indicates that, based on the bridge given in Theorem 1, the fault features can be transformed between each other. Also, (53) can be used as an auxiliary evidential statement to illustrate the validity of Theorems 2 and 3 because of

$$\begin{aligned} \Sigma_{\mathbf{r}_s^{\text{unsp}}} &= \mathbb{E}(\mathbf{e}_s(k) \mathbf{e}_s^T(k)), \\ \Sigma_{\mathbf{r}_s^{\text{sp}}} &= \mathcal{P}_{\text{sp}/\text{unsp}} \Sigma_{\mathbf{r}_s^{\text{unsp}}} \mathcal{P}_{\text{sp}/\text{unsp}}^T, \end{aligned} \quad (54)$$

such that

$$\begin{aligned} \mathbb{E}(T^2(\mathbf{r}_s^{\text{unsp}}(k))) &= \mathbb{E}(T^2(\mathbf{r}_s^{\text{sp}}(k))), \\ \mathbb{E}(T^2(\mathbf{r}_s^{\text{unsp},f}(k))) &= \mathbb{E}(T^2(\mathbf{r}_s^{\text{sp},f}(k))). \end{aligned} \quad (55)$$

2) *Geometric interpretation:* Fig. 2 provides the geometric descriptions of both unsupervised and supervised learning-aided parameter identification, where the bridge, $\mathcal{P}_{\text{unsp}/\text{sp}}$ and $\mathcal{P}_{\text{sp}/\text{unsp}}$, is highlighted by the two red curves. It is worth mentioning that $\mathbf{y}^{\text{unsp},*}$ and $\mathbf{y}^{\text{sp},*}$ reconstruct the dynamic behaviors of nonlinear systems by carrying the maximum variation information (i.e., the largest uncertainty). They are suitable for nonlinear parameter identification. At the same time, $\mathcal{E}^{\text{unsp},*}$ and $\mathcal{E}^{\text{sp},*}$ with the minimum uncertainty are the best choices for fault diagnosis [16].



(a) unsupervised learning (b) supervised learning

Fig. 2: Geometric descriptions of the suspected and residual spaces in system identification.

Fig. 3 sketches the changes caused by $\mathbf{f}_a$ and $\mathbf{f}_s$ in residual spaces, together with the same bridge in both normal and faulty scenarios.

3) *A generalized version:* In Theorems 1 to 3, $\mathcal{M}_{\text{euc}}$ is chosen as the metric to measure the difference not only between $\mathbf{y}_s$ and $\mathbf{y}_s^{\text{unsp}}$ when utilizing unsupervised learning to identify $\mathcal{K}_{s_p+s}^{\text{unsp}}$, but also between $\mathbf{y}_s$ and $\mathbf{y}_s^{\text{sp}}$ when utilizing supervised learning to estimate $\mathcal{K}_{s_p+s}^{\text{sp}}$. For the fault diagnosis purpose, the same performance can be obtained by using $\mathcal{M}_{\text{cor}}$ defined in (4a)-(4b). As pointed in [23] and [25], an IFD design relying on the $\mathcal{M}_{\text{cor}}$ metric is a generalized version

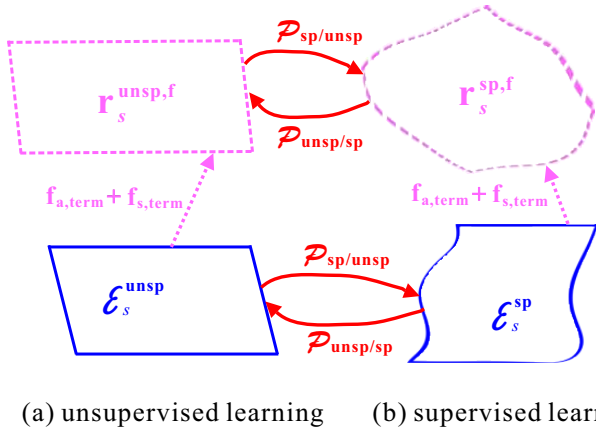


Fig. 3: Geometric descriptions of the residual spaces in both normal and faulty scenarios.

that can provide multiple optimal solutions to both $\mathcal{K}_{s_p+s}^{\text{unsp}}$ and $\mathcal{K}_{s_p+s}^{\text{sp}}$.

In order to distinguish it from the solution based on $\mathcal{M}_{\text{euc}}$, the notations in the generalized version are marked by “ $\underline{\cdot}$ ”. Therefore, the objective functions of two generalized versions, respectively corresponding to unsupervised and supervised learning, can be formulated as follows.

$$\begin{aligned} & \min \mathcal{M}_{\text{cor}}(\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{unsp}}), \\ & = \min k_{\mathbf{y}_s} - \text{Tr} \left(\left| \Sigma_{\mathbf{y}_s}^{-1/2} \Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{unsp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1/2} \right| \right), \quad (56) \\ & \iff \min k_{\mathbf{y}_s} - \text{Tr} \left(\Sigma^{\text{unsp}, T} \Sigma^{\text{unsp}} \right), \end{aligned}$$

and

$$\begin{aligned} & \min \mathcal{M}_{\text{cor}}(\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{sp}}), \\ & = \min k_{\mathbf{y}_s} - \text{Tr} \left(\left| \Sigma_{\mathbf{y}_s}^{-1/2} \Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{sp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{sp}}}^{-1/2} \right| \right), \quad (57) \\ & \iff \min k_{\mathbf{y}_s} - \text{Tr} \left(\Sigma^{\text{sp}, T} \Sigma^{\text{sp}} \right), \end{aligned}$$

where Σ^{unsp} and Σ^{sp} are obtained via the following singular value decompositions:

$$\begin{aligned} \Sigma_{\mathbf{y}_s}^{-1/2} \Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{unsp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1/2} &= \underline{\Gamma}^{\text{unsp}} \underline{\Sigma}^{\text{unsp}} \underline{\Upsilon}^{\text{unsp}, T}, \\ \Sigma_{\mathbf{y}_s}^{-1/2} \Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{sp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{sp}}}^{-1/2} &= \underline{\Gamma}^{\text{sp}} \underline{\Sigma}^{\text{sp}} \underline{\Upsilon}^{\text{sp}, T}. \end{aligned} \quad (58)$$

Now we define two sets of linear mappings as follows.

$$\underline{\mathbf{G}}_{\mathbf{y}_s}^{\text{unsp}} = \underline{\Gamma}^{\text{unsp}, T} \Sigma_{\mathbf{y}_s}^{-1/2}, \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}} = \underline{\Upsilon}^{\text{unsp}, T} \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1/2}, \quad (59a)$$

$$\underline{\mathbf{G}}_{\mathbf{y}_s}^{\text{sp}} = \underline{\Gamma}^{\text{sp}, T} \Sigma_{\mathbf{y}_s}^{-1/2}, \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{sp}}}^{\text{sp}} = \underline{\Upsilon}^{\text{sp}, T} \Sigma_{\underline{\mathbf{y}}_s^{\text{sp}}}^{-1/2}. \quad (59b)$$

Then the two generalized residual generators can be constructed according to

$$\underline{\mathbf{r}}_s^{\text{unsp}} = \underline{\mathbf{G}}_{\mathbf{y}_s}^{\text{unsp}} \mathbf{y}_s - \underline{\Sigma}^{\text{unsp}} \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}} \underline{\mathbf{y}}_s^{\text{unsp}}, \quad (60a)$$

$$\underline{\mathbf{r}}_s^{\text{sp}} = \underline{\mathbf{G}}_{\mathbf{y}_s}^{\text{sp}} \mathbf{y}_s - \underline{\Sigma}^{\text{sp}} \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{sp}}}^{\text{sp}} \underline{\mathbf{y}}_s^{\text{sp}}. \quad (60b)$$

It can be verified that

$$\Sigma_{\underline{\mathbf{r}}_s^{\text{unsp}}} = \mathbf{I} - \underline{\Sigma}^{\text{unsp}} \underline{\Sigma}^{\text{unsp}, T}, \Sigma_{\underline{\mathbf{r}}_s^{\text{sp}}} = \mathbf{I} - \underline{\Sigma}^{\text{sp}} \underline{\Sigma}^{\text{sp}, T} \quad (61)$$

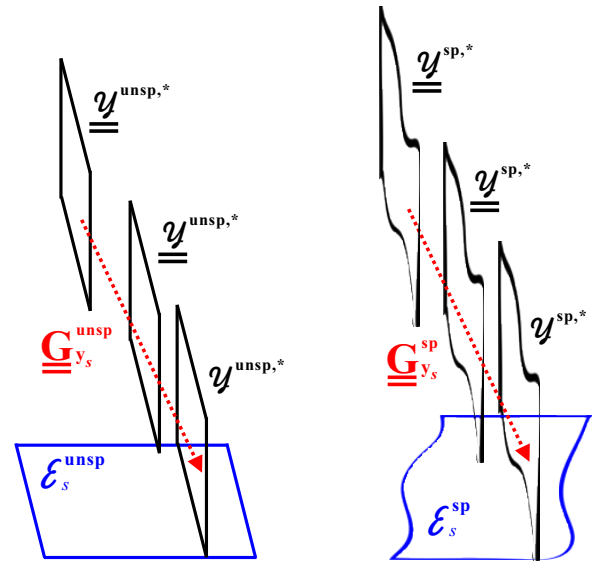


Fig. 4: More general versions of suspected and residual spaces.

Furthermore, by substituting (59a) into (60a), one can obtain

$$\begin{aligned} & \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}, -1} \underline{\mathbf{r}}_s^{\text{unsp}} \\ &= \underline{\mathbf{y}}_s - \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}, -1} \underline{\Sigma}^{\text{unsp}} \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}} \underline{\mathbf{y}}_s^{\text{unsp}} \\ &= \underline{\mathbf{y}}_s - \Sigma_{\mathbf{y}_s}^{1/2} (\underline{\Gamma}^{\text{unsp}} \underline{\Sigma}^{\text{unsp}} \underline{\Upsilon}^{\text{unsp}, T}) \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1/2} \underline{\mathbf{y}}_s^{\text{unsp}} \\ &= \underline{\mathbf{y}}_s - \Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{unsp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1} \underline{\mathbf{y}}_s^{\text{unsp}} \end{aligned} \quad (62)$$

because $\underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}}$ has full rank. Define

$$\underline{\mathbf{y}}_s^{\text{unsp}} = \underline{\mathbf{U}}_{\mathbf{y}} \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} \quad (63)$$

where its optimal choices, $\underline{\mathbf{y}}_s^{\text{unsp},*}$ and $\underline{\mathbf{U}}_{\mathbf{y}}^*$, are obtained according to (56). It must be pointed out that $\underline{\mathbf{U}}_{\mathbf{y}}^*$ and its corresponding output $\underline{\mathbf{y}}_s^{\text{unsp},*}$ are not unique, because the objective functions are defined by $\mathcal{M}_{\text{euc}}$. Then (62) can be further written as

$$\begin{aligned} & \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}, -1} \underline{\mathbf{r}}_s^{\text{unsp}} \\ &= \underline{\mathbf{y}}_s - \underbrace{\Sigma_{\mathbf{y}_s, \underline{\mathbf{y}}_s^{\text{unsp}}} \Sigma_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{-1} \underline{\mathbf{U}}_{\mathbf{y}}^*}_{\underline{\mathbf{U}}_{\mathbf{y}}^*} \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} \end{aligned} \quad (64)$$

$$\begin{aligned} &= \mathbf{e}_s = \mathbf{r}_s^{\text{unsp}} \\ &\iff \underline{\mathbf{r}}_s^{\text{unsp}} = \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{unsp}}}^{\text{unsp}} \underline{\mathbf{r}}_s^{\text{unsp}}. \end{aligned}$$

Similarly, the following relationship holds for supervised learning-based IFD approaches:

$$\underline{\mathbf{r}}_s^{\text{sp}} = \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s^{\text{sp}}}^{\text{sp}} \underline{\mathbf{r}}_s^{\text{sp}}. \quad (65)$$

Eventually, the following two relationships hold:

$$\begin{aligned} T^2(\underline{\mathbf{r}}_s^{\text{unsp}}(k)) &= \underline{\mathbf{r}}_s^{\text{unsp},T}(k) \underline{\Sigma}_{\underline{\mathbf{r}}_s^{\text{unsp}}}^{-1} \underline{\mathbf{r}}_s^{\text{unsp}}(k) \\ &= (\underline{\mathbf{G}}_{\underline{\mathbf{y}}_s}^{\text{unsp}} \underline{\mathbf{r}}_s^{\text{unsp}})^T \underline{\Sigma}_{\underline{\mathbf{r}}_s^{\text{unsp}}}^{-1} \underline{\mathbf{G}}_{\underline{\mathbf{y}}_s}^{\text{unsp}} \underline{\mathbf{r}}_s^{\text{unsp}} = T^2(\underline{\mathbf{r}}_s^{\text{unsp}}(k)) \end{aligned} \quad (66)$$

and

$$T^2(\underline{\mathbf{r}}_s^{\text{sp}}(k)) = T^2(\underline{\mathbf{r}}_s^{\text{sp}}(k)). \quad (67)$$

Combining (67) with (48) obtains

$$\begin{aligned} T^2(\underline{\mathbf{r}}_s^{\text{sp}}(k)) &= T^2(\underline{\mathbf{r}}_s^{\text{unsp}}(k)) = \\ T^2(\underline{\mathbf{r}}_s^{\text{sp}}(k)) &= T^2(\underline{\mathbf{r}}_s^{\text{unsp}}(k)), \end{aligned} \quad (68)$$

for normal operations, and

$$\begin{aligned} T^2(\underline{\mathbf{r}}_s^{\text{sp},f}(k)) &= T^2(\underline{\mathbf{r}}_s^{\text{unsp},f}(k)) = \\ T^2(\underline{\mathbf{r}}_s^{\text{sp},f}(k)) &= T^2(\underline{\mathbf{r}}_s^{\text{unsp},f}(k)), \end{aligned} \quad (69)$$

for faulty operations.

Following Theorems 2 and 3, T^2 test statistic defined on the two generalized residual generators also have the optimal fault-detection power. The readers can refer to [23] for the more rigorous analysis of the generalized version.

Remark 6. We can find that the generalized versions of residual generators also deliver the least-square estimation of $\underline{\mathbf{y}}_s$, because both $\underline{\mathbf{G}}_{\underline{\mathbf{y}}_s}^{\text{unsp}}$ and $\underline{\mathbf{G}}_{\underline{\mathbf{y}}_s}^{\text{sp}}$ play a role as normalization by limiting their covariance matrices to (61). ∇

In order to have an insightful observation, Fig. 4 depicts multiple solutions to the suspected spaces, where Figs. 4(a) and (b) respectively correspond to unsupervised and supervised learning strategies.

IV. IFD DESIGNS AND IMPLEMENTATIONS: FROM UNSUPERVISED TO SUPERVISED NEURAL NETWORKS

A study of the existing approaches reveals that both the unsupervised and supervised machine learning techniques can be employed to enhance the IFD flexibility against system nonlinearity. For example, neural networks have received an increasing attention for the designs of IFD. Directly motivated by Theorems 1 to 3, the fully connected neural networks are used in this section for the fault diagnosis purpose.

A. Unit-time Delay Operators

A recurrent neural network is capable of describing the dynamic behaviors of nonlinear systems, thanks to the unit-time delays attached in neurons [30]. As the system order increases, more delay units are necessary to fit higher order dynamics, resulting in heavy computations together with problems associated with vanishing and exploding gradients [31]. To avoid such a situation and improve computation efficiency, the unit-time delay operator, denoted as z^{-1} , is used in the data pre-processing stage.

As shown in Fig. 5(a), $(s_p + s + 1)$ unit-time delay operators are defined on both $\mathbf{u}(k + s)$ and $\mathbf{y}(k + s)$ to obtain the inputs

of an unsupervised neural network, i.e.,

$$\begin{aligned} \mathbf{u}(k + s - 1) &= z^{-1} \mathbf{u}(k + s), \quad \dots, \\ \mathbf{u}(k - s_p - 1) &= z^{-(s_p + s + 1)} \mathbf{u}(k + s), \\ \mathbf{y}(k + s - 1) &= z^{-1} \mathbf{y}(k + s), \quad \dots, \\ \mathbf{y}(k - s_p - 1) &= z^{-(s_p + s + 1)} \mathbf{y}(k + s). \end{aligned} \quad (70)$$

Similarly, $(s_p + s + 1)$ and $(s_p + 1)$ unit-time delay operators are defined on $\mathbf{u}(k + s)$ and $\mathbf{y}(k)$, respectively. The inputs of a supervised neural network given in Fig. 5(b) are $[\mathbf{u}_p^T \ \mathbf{u}_s^T]$ and

$$\begin{aligned} \mathbf{y}(k - 1) &= z^{-1} \mathbf{y}(k), \quad \dots, \\ \mathbf{y}(k - s_p - 1) &= z^{-(s_p + 1)} \mathbf{y}(k). \end{aligned} \quad (71)$$

In addition, the unit-time delays are also used to obtain the reference outputs of both the unsupervised and supervised neural networks.

B. Unsupervised IFD Approaches

Similar to (32a), the loss function L of the unsupervised neural network can be defined by

$$L^{\text{unsp}} = \left\| \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} - \underbrace{\mathcal{H}^{\text{unsp}}(\Theta_z \ \Theta_u \ \Theta_y)}_{\mathcal{H}^{\text{unsp}}(\Theta)} \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} \right\|_2^2, \quad (72)$$

where $\mathcal{H}^{\text{unsp}}(\Theta)$ is the architecture of the neural network with the hyper parameter Θ , and the subscripts of Θ signify their corresponding outputs. Therefore, the optimal Θ can be obtained by minimizing L^{unsp} , i.e.,

$$\Theta^* = (\Theta_z^* \ \Theta_u^* \ \Theta_y^*) = \arg \min_{\Theta} L^{\text{unsp}}, \quad (73)$$

based on which the residual signal using $\mathcal{H}^{\text{unsp}}(\Theta^*)$ is

$$\mathbf{r}_s^{\text{unsp}}(k) = \mathbf{y}_s - \mathcal{H}^{\text{unsp}}(\Theta_y^*) \circ \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} = (27a). \quad (74)$$

In Fig. 5 (a), $\mathbf{r}_s^{\text{unsp}}(k)$ is highlighted in red and “ $\times$ ” refers to the reconstruction errors that may be neglected.

Combining with Fig. 5 (a), the implementation procedures of an unsupervised neural network-based IFD approach are summarized in Algorithms 1 and 2.

C. Supervised IFD Approaches

Considering a supervised neural network $\mathcal{H}^{\text{sp}}(\Theta)$, its loss function is formulated as follows

$$L^{\text{sp}} = \left\| \mathbf{y}_s - \mathcal{H}^{\text{sp}}(\Theta) \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \end{bmatrix} \right\|_2^2, \quad (76)$$

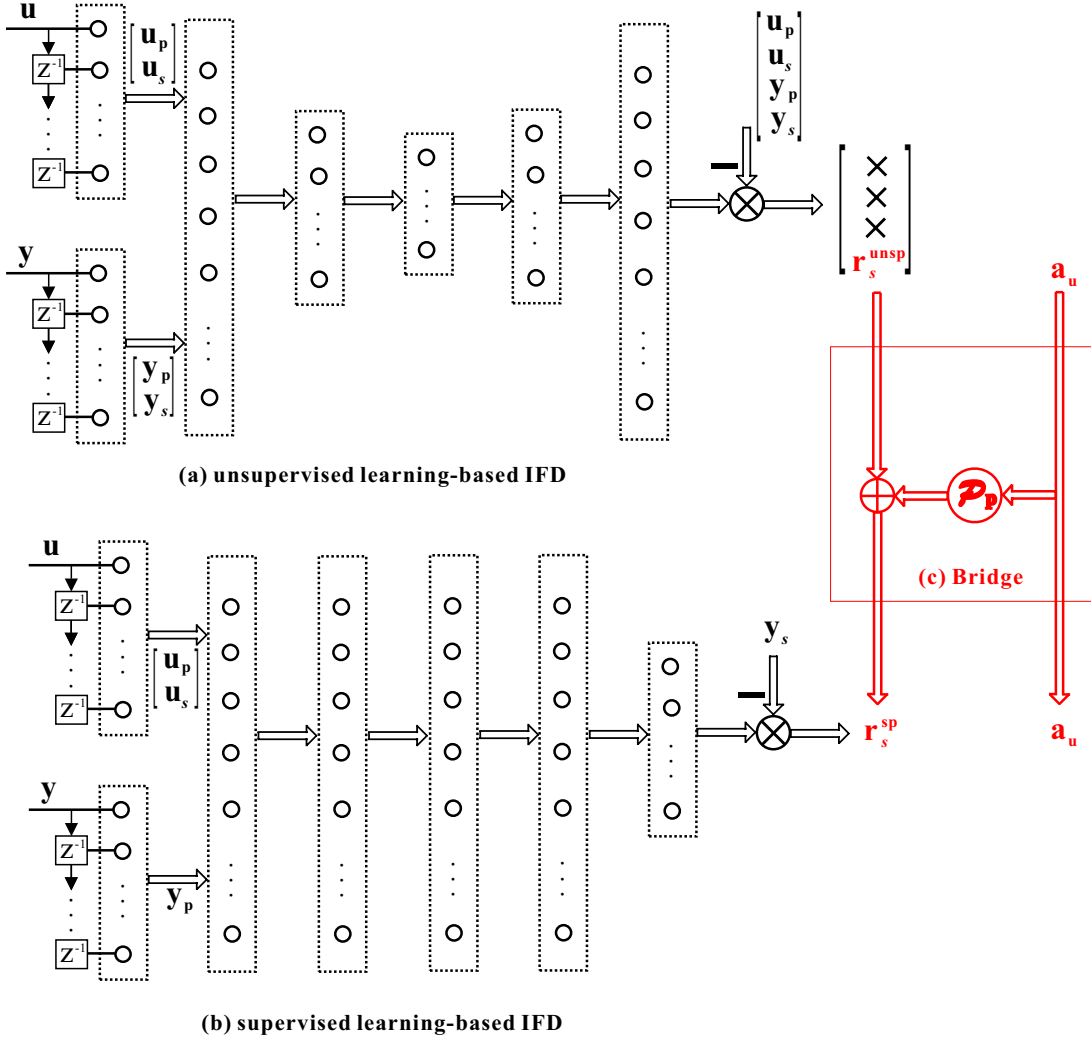


Fig. 5: Schematic of IFD frameworks with multiple unit-time delays, where “ a_u ” represents an auxiliary variable, (a) is designed based an unsupervised neural network, (b) is implemented based a supervised neural network, and (c) highlighted in red is the bridge using an invertible neural network.

Algorithm 1 Off-line Learning: Unsupervised neural network-based IFD approaches

- 1: Collect system measurements $\mathbf{u}$ and $\mathbf{y}$;
 - 2: Pre-process $\mathbf{u}$ and $\mathbf{y}$ by using the delay units according to (70) to obtain $\mathbf{z}_p$, $\mathbf{u}_s$ and $\mathbf{y}_s$;
 - 3: Construct an unsupervised neural network whose architecture is $\mathcal{H}^{\text{unsp}}(\Theta)$ and loss function is defined in (72);
 - 4: Update Θ according to (73) and obtain Θ_y^* ;
 - 5: Generate the residual signal $\mathbf{r}_s^{\text{unsp}}$ via (74);
 - 6: Determine the threshold J_{th} for T^2 based on (25);
 - 7: (If necessary) identify the fault features based on system knowledge or using faulty data.
-

when $\mathcal{M}_{\text{euc}}^2(\mathbf{y}_s, \mathbf{y}_s^{\text{sp}})$ is adopted. Then the optimal hyper parameter can be obtained according to

$$\Theta^* = \arg \min_{\Theta} L^{\text{sp}}, \quad (77)$$

Algorithm 2 Online Fault Diagnosis: Unsupervised neural network-based IFD approaches

- 1: Read real-time data and employ unit-time delay operators;
 - 2: Calculate the online residual signal according to (74);
 - 3: Compute the test statistic via (24);
 - 4: Make an FD decision according to

$$T^2(\mathbf{r}_s^{\text{unsp}}(\text{online})) - J_{th} < 0 \implies \text{Fault-free};$$

$$\text{Otherwise} \implies \text{Faulty}; \quad (75)$$
 - 5: (If necessary) diagnose the fault by using $\mathbf{r}_s^{\text{unsp}}(\text{online})$.
-

which allows for constructing the residual signal as follows.

$$\mathbf{r}_s^{\text{sp}}(k) = \mathbf{y}_s - \mathcal{H}^{\text{sp}}(\Theta^*) \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \end{bmatrix} = (27b). \quad (78)$$

According to the analysis above and the schematic presented in Fig. 5(b), the complete designs and implementation proce-

dures are given in Algorithms 3 and 4.

Algorithm 3 Off-line Learning: Supervised neural network-based IFD approaches

- 1: Collect system measurements $\mathbf{u}$ and $\mathbf{y}$;
 - 2: Pre-process $\mathbf{u}$ and $\mathbf{y}$ by using the delay units according to (70) and (71) to obtain $\mathbf{z}_p$ and $\mathbf{u}_s$;
 - 3: Construct a supervised neural network whose architecture is $\mathcal{H}^{sp}(\Theta)$ and loss function is defined in (76);
 - 4: Update Θ via (77) and obtain Θ^* ;
 - 5: Generate the residual signal $\mathbf{r}_s^{sp}$ via (78);
 - 6: Determine the threshold J_{th} for T^2 based on (25);
 - 7: (If necessary) identify the fault features based on system knowledge or using faulty data.
-

Algorithm 4 Online Fault Diagnosis: Supervised neural network-based IFD approaches

- 1: Read real-time data and employ unit-time delay operators;
- 2: Calculate the online residual signal according to (78);
- 3: Compute the test statistic via (24);
- 4: Make an FD decision according to

$$T^2(\mathbf{r}_s^{sp}(\text{online})) - J_{th} < 0 \implies \text{Fault-free;} \quad (79)$$

$$\text{Otherwise} \implies \text{Faulty;} \quad (79)$$

- 5: (If necessary) diagnose the fault by using $\mathbf{r}_s^{sp}(\text{online})$.
-

D. An Invertible Neural Network-aided Bridge

In order to build the bridge between unsupervised and supervised neural networks-based IFD approaches, we introduce an auxiliary variable $\mathbf{a}_u$. The mapping generated by $\mathcal{H}(\Theta_{\text{ext}})$ in the red square of Fig. 5(c) is

$$\begin{bmatrix} \mathbf{r}_s^{sp} \\ \mathbf{a}_u \end{bmatrix} = \underbrace{\begin{pmatrix} \mathbf{I} & \mathcal{H}(\Theta_p) \\ 0 & \mathbf{I} \end{pmatrix}}_{\mathcal{H}(\Theta_{\text{ext}})} \begin{bmatrix} \mathbf{r}_s^{\text{unsp}} \\ \mathbf{a}_u \end{bmatrix} \quad (80)$$

where the loss function of $\mathcal{H}(\Theta_{\text{ext}})$ is chosen as

$$L^{\text{ext}} = \left\| \begin{bmatrix} \mathbf{r}_s^{sp} \\ \mathbf{a}_u \end{bmatrix} - \mathcal{H}(\Theta_{\text{ext}}) \begin{bmatrix} \mathbf{r}_s^{\text{unsp}} \\ \mathbf{a}_u \end{bmatrix} \right\|_2^2. \quad (81)$$

Without loss of generality, we can choose $\mathbf{a}_u = \mathbf{r}_s^{\text{unsp}}$ to prove the existence of the bridge. Then minimizing L^{ext} obtains

$$\Theta_{\text{ext}}^* = \arg \min_{\Theta_{\text{ext}}} L^{\text{ext}} \iff \mathcal{H}(\Theta_p^*) = \mathcal{P}_p \quad (82)$$

as shown in Fig. 5(c). It indicates that

$$\mathbf{r}_s^{sp} = (\mathcal{P}_p + \mathbf{I})\mathbf{r}_s^{\text{unsp}}, \mathbf{r}_s^{\text{unsp}} = \mathbf{r}_s^{\text{unsp}}, \quad (83)$$

and

$$\begin{aligned} \mathbf{r}_s^{\text{unsp}} &= \mathbf{r}_s^{sp} - \mathcal{P}_p \mathbf{r}_s^{\text{unsp}} = \mathbf{r}_s^{sp} - \mathcal{H}(\Theta_p^*) \mathbf{r}_s^{\text{unsp}} \\ &= \mathbf{r}_s^{sp} - \mathcal{H}(\Theta_p^*) \mathbf{a}_u, \\ \mathbf{a}_u &= \mathbf{a}_u. \end{aligned} \quad (84)$$

It is interesting to see that, even though the inverse operation is not used, we can obtain the bridge by

$$\begin{bmatrix} \mathbf{r}_s^{sp} \\ \mathbf{a}_u \end{bmatrix} = \begin{pmatrix} \mathbf{I} & \mathcal{P}_p \\ 0 & \mathbf{I} \end{pmatrix} \begin{bmatrix} \mathbf{r}_s^{\text{unsp}} \\ \mathbf{a}_u \end{bmatrix}, \quad (85)$$

and

$$\begin{bmatrix} \mathbf{r}_s^{\text{unsp}} \\ \mathbf{a}_u \end{bmatrix} = \begin{pmatrix} \mathbf{I} & -\mathcal{P}_p \\ 0 & \mathbf{I} \end{pmatrix} \begin{bmatrix} \mathbf{r}_s^{sp} \\ \mathbf{a}_u \end{bmatrix}. \quad (86)$$

In addition, the Jacobian matrix $\mathbf{J}$ of $\mathcal{H}(\Theta_{\text{ext}})$ is

$$\mathbf{J} = \begin{pmatrix} \mathbf{I} & \mathcal{P}_p \\ 0 & \mathbf{I} \end{pmatrix}, \quad (87)$$

whose determinant is always positive, i.e.,

$$|\mathbf{J}| = \left| \begin{pmatrix} \mathbf{I} & \mathcal{P}_p \\ 0 & \mathbf{I} \end{pmatrix} \right| = 1. \quad (88)$$

It means that $\mathcal{H}(\Theta_{\text{ext}})$ and its component $\mathcal{P}_{sp/unsp}$ are globally invertible.

Also, (85), (86), and (88) can be used as an auxiliary evidence to illustrate the validity of Theorem 1.

Remark 7. The solutions to the bridge depicted in Fig. 5(c) are not unique. For example, a “reversible residual network” given in [32] can also be directly employed to obtain an one-to-one mapping between $\mathbf{r}_s^{\text{unsp}}$ and $\mathbf{r}_s^{sp}$. ∇

V. CONCLUSIONS

Thanks to the advanced machine learning techniques, there appears to have room for further development of nonlinear IFD. System data can provide us with information to reveal physical principles, making more informative inferences and decisions feasible. As the interpretability of these advanced methods is fundamental, the interplay between physical principles (such as a mathematical system description whose parameters may be unknown) and data becomes more critical.

This perspective article has developed three theorems related to IFD and related parameter identification for nonlinear dynamic systems. In order to obtain a unified form, Theorem 1 builds a bridge between unsupervised and supervised learning-based residual generators. Theorems 2 and 3, respectively corresponding to unsupervised and supervised learning, develop the IFD structures and illustrate their optimal performance for fault detection. With the aid of three different kinds of neural networks, Section IV details the specific designs and implementations of Theorems 1 to 3. This work lays a foundation for the further development of explainable IFD.

The perspective article ends with our opinions, expected challenges, and future research opportunities as follows.

- Not all IFD approaches are satisfactory from the viewpoint of fault-diagnosis performance. Our work provides researchers with some instructive guidance on designing more effective IFD algorithms, including both unsupervised and supervised learning-based schemes.
- As shown in Fig. 5 and Theorem 2, not all results obtained through unsupervised learning are useful for fault diagnosis. The conclusion is also true and easier to explain for the case of linear dynamic systems. For example, the purpose of singular value decomposition [33] and QR [5] used in linear approaches is not for dimension reduction, although they can do so. Keep in mind that the ultimate goal of unsupervised learning adopted in IFD approaches is always to minimize the reconstruction error of system outputs for both the linear and nonlinear systems.
- The initial step to apply the results obtained through this work must be to show the existence of the bridge $\mathcal{P}_{\text{sp/unsup}}$ and $\mathcal{P}_{\text{unsup/sp}}$ given in Theorem 1. Then a series of invertible machine learning tools, such as invertible neural networks in [34], can be employed to build this bridge. The process should not be reversed.
- An unmentioned condition in this study for nonlinear IFD approaches is that (6) is output reconstructible [35]. In fact, the condition is weak and easy to satisfy for our proposed IFD frameworks, because a multi-layer neural network can approximate a continuous function (such as $\phi(\cdot)$ and $v(\cdot)$ in Fig. 1) arbitrarily accurate [36].
- The data-based model (12) is obtained from (6) with the assumption of additivity given in (11a)-(11b). In practical applications, many classes of nonlinear systems satisfy (11a)-(11b) such as bounded systems, Hammerstein systems, and affine systems.
- It is a fact that nonlinear IFD approaches and the associated thresholds show dependence on $\mathbf{u}$ and $\mathbf{x}(0)$ [13]. In order to obtain a reasonable data set for training, a uniformly distributed $\mathbf{u}$ over different operation points is recommended so that a persistent excitation of the global system nonlinearities can be achieved.
- In addition to the bridge developed in this work, other aspects of system identification and IFD approaches deserve more investigations, such as (a) How many neurons should be used when utilizing a neural network? (b) How to determine the nonlinearity degree to approximate an unknown dynamic system without the problem of overfitting or underfitting? (c) What is the minimum size of training data for obtaining a desired performance? and so on.

APPENDIX A THE PROOF OF THEOREM 1

Given an arbitrary unsupervised learning approach whose notation is $\mathcal{U}$, it generates the nonlinear mapping (28) with

the objective function such as defined via $\mathcal{M}_{\text{euc}}^2$:

$$\min \left\| \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} - \begin{bmatrix} \hat{\mathbf{z}}_p \\ \hat{\mathbf{u}}_s \\ \hat{\mathbf{y}}_s \end{bmatrix} \right\|_2^2 \Rightarrow \min \|\mathbf{y}_s - \hat{\mathbf{y}}_s\|_2^2 = \min \|\mathbf{e}_s\|_2^2. \quad (\text{A.1})$$

In (28), the input of $\mathcal{U}_z$, $\mathcal{U}_u$, and $\mathcal{U}_y$ is $[\mathbf{z}_p^T \mathbf{u}_s^T \mathbf{y}_s^T]^T$, and the subscripts correspond to the output variables.

Three steps are necessary to complete the proof.

Step 1): The generation of innovation error $\mathbf{e}_s$ using unsupervised learning. Based on (16a) and (16b), one can obtain

$$\begin{aligned} \mathbf{y}(k) &= \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \hat{\mathbf{x}}(k) + \mathbf{v}_{\hat{\mathbf{y}}\mathbf{u}} \mathbf{u}(k) + \mathbf{e}(k) \\ \mathbf{y}(k+1) &= \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \hat{\mathbf{x}}(k+1) + \mathbf{v}_{\hat{\mathbf{y}}\mathbf{u}} \mathbf{u}(k+1) + \mathbf{e}(k+1) \\ &= \mathbf{v}_{\hat{\mathbf{y}}\mathbf{u}} \mathbf{u}(k+1) + \mathbf{e}(k+1) + \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{A}_\ell \mathbf{y}(k) + \\ &\quad \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} (\underbrace{\phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}} - \ell \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}}}_{\mathcal{A}_\ell}) \hat{\mathbf{x}}(k) + \\ &\quad \mathbf{v}_{\hat{\mathbf{y}}\hat{\mathbf{x}}} (\underbrace{\phi_{\hat{\mathbf{x}}\mathbf{u}} - \ell \mathbf{v}_{\hat{\mathbf{y}}\mathbf{u}}}_{\mathcal{B}_\ell}) \mathbf{u}(k) \\ &\dots \end{aligned} \quad (\text{A.2})$$

Then, $\mathbf{e}_s(k)$ can be described by

$$\mathbf{e}_s(k) = \mathbf{y}_s(k) - \mathbf{\Upsilon}_y \mathbf{y}_s(k) - \mathbf{\Upsilon}_x \hat{\mathbf{x}}(k) - \mathbf{\Upsilon}_u \mathbf{u}_s(k) \quad (\text{A.3})$$

where $\mathbf{\Upsilon}_y$, $\mathbf{\Upsilon}_x$, and $\mathbf{\Upsilon}_u$ are the nonlinear composite operators:

$$\mathbf{\Upsilon}_y = \begin{pmatrix} 0 & 0 & \dots & 0 \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \ell & 0 & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{A}_\ell^{s-1} \ell & \dots & v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \ell & 0 \end{pmatrix}, \quad (\text{A.4})$$

$$\mathbf{\Upsilon}_u = \begin{pmatrix} v_{\hat{\mathbf{y}}\mathbf{u}} & 0 & \dots & 0 \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{B}_\ell & v_{\hat{\mathbf{y}}\mathbf{u}} & \dots & 0 \\ \vdots & \ddots & \ddots & \vdots \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{A}_\ell^{s-1} \mathcal{B}_\ell & \dots & v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{B}_\ell & v_{\hat{\mathbf{y}}\mathbf{u}} \end{pmatrix}, \quad (\text{A.5})$$

$$\mathbf{\Upsilon}_x = \begin{pmatrix} v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \\ \vdots \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \mathcal{A}_\ell^s \end{pmatrix}. \quad (\text{A.6})$$

Through some mathematical manipulations, (19) can be rewritten as

$$\mathbf{y}_s(k) = \mathbf{\Upsilon}_x \hat{\mathbf{x}}(k) + \mathbf{\Upsilon}_u \mathbf{u}_s(k) + \mathbf{\Upsilon}_e \mathbf{e}_s(k) \quad (\text{A.7})$$

which can be achieved by using supervised learning approaches, as shown in (19), i.e.,

$$\hat{\mathbf{y}}_s(k) = \underbrace{(\mathbf{\Upsilon}_x \mathcal{L}_p \quad \mathbf{\Upsilon}_u)}_{\mathcal{S} := (\mathcal{S}_z \quad \mathcal{S}_u)} \begin{bmatrix} \mathbf{z}_p(k) \\ \mathbf{u}_s(k) \end{bmatrix} \quad (\text{A.8})$$

where $\mathcal{S}$ is the nonlinear composite operator generated by supervised learning.

Step 2): A unified description of the transformation between $\mathbf{r}_s^{\text{unsp}}$ and $\mathbf{r}_s^{\text{sp}}$. Based on (A.3) and (A.8), two kinds of residual generators can be described by

$$\mathbf{r}_s^{\text{unsp}} = \mathbf{y}_s - \mathcal{U}_y \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \\ \mathbf{y}_s \end{bmatrix} = \mathbf{e}_s, \quad (\text{A.9a})$$

$$\mathbf{r}_s^{\text{sp}} = \mathbf{y}_s - (\mathcal{S}_z \ \mathcal{S}_u) \begin{bmatrix} \mathbf{z}_p \\ \mathbf{u}_s \end{bmatrix} = \Upsilon_e \mathbf{e}_s \quad (\text{A.9b})$$

Step 3): The existence of $\mathcal{P}_{\text{sp/unsp}}^{-1}$. In fact, $\mathbf{r}_s^{\text{unsp}}$ in (A.9a) and $\mathbf{r}_s^{\text{sp}}$ in (A.9b) have shown

$$\mathbf{r}_s^{\text{sp}} = \mathcal{P}_{\text{sp/unsp}} \mathbf{r}_s^{\text{unsp}}, \quad \mathcal{P}_{\text{sp/unsp}} := \Upsilon_e. \quad (\text{A.10})$$

For the sake of simplicity, we rewrite (A.10) as

$$\mathbf{r}_s^{\text{sp}} = \mathbf{r}_s^{\text{unsp}} + \mathcal{P}_p \mathbf{r}_s^{\text{unsp}} \quad (\text{A.11})$$

where $\mathcal{P}_p$ is the component of $\mathcal{P}_{\text{sp/unsp}}$, i.e.,

$$\mathcal{P}_p = \begin{pmatrix} \mathbf{0} & \cdots & \mathbf{0} \\ \vdots & \ddots & \vdots \\ v_{\hat{\mathbf{y}}\hat{\mathbf{x}}} \phi_{\hat{\mathbf{x}}\hat{\mathbf{x}}}^{s-1} \ell_{\hat{\mathbf{x}}\mathbf{e}} & \cdots & \mathbf{0} \end{pmatrix}. \quad (\text{A.12})$$

Consider a nonlinear operator $\mathcal{P}_{\text{unsp/sp}}$ and the Lipschitz constant of $\mathcal{P}_p$ is $\text{Lip} < 1$. By defining a sequence $\{\mathbf{e}_{s,j}\} \in \mathbf{e}_s$, one can obtain the following innovation:

$$\begin{aligned} \mathcal{P}_{\text{unsp/sp}} : \mathbf{e}_{s,j+1} &= \Upsilon_e \mathbf{e}_s - \mathcal{P}_p \mathbf{e}_{s,j} \\ \implies \lim_{j \rightarrow \infty} \mathbf{e}_{s,j} &= \mathcal{P}_{\text{unsp/sp}} \circ \Upsilon_e \mathbf{e}_s. \end{aligned} \quad (\text{A.13})$$

Given a positive integer n , we have

$$\begin{aligned} \|\mathbf{e}_{s,j+n} - \mathbf{e}_{s,j}\|_2 &\leq \|\mathbf{e}_{s,j+n} - \mathbf{e}_{s,j+n-1}\|_2 + \cdots \\ &\quad \|\mathbf{e}_{s,j+1} - \mathbf{e}_{s,j}\|_2 \\ &= \|\mathcal{P}_p(\mathbf{e}_{s,j+n-1} - \mathbf{e}_{s,j+n-2})\|_2 + \cdots + \\ &\quad \|\mathcal{P}_p(\mathbf{e}_{s,j} - \mathbf{e}_{s,j-1})\|_2 \\ &\leq (\text{Lip}^{j+n-1} + \cdots + \text{Lip}^j) \|\mathbf{e}_{s,1} - \mathbf{e}_{s,0}\|_2 \\ &= \frac{1 - \text{Lip}^{n-1}}{1 - \text{Lip}} \text{Lip}^j \|\mathbf{e}_{s,1} - \mathbf{e}_{s,0}\|_2 \\ &\leq \frac{\text{Lip}^j}{1 - \text{Lip}} \|\mathbf{e}_{s,1} - \mathbf{e}_{s,0}\|_2. \end{aligned} \quad (\text{A.14})$$

It indicates that given an arbitrary small ε , there always exists a n such that

$$\|\mathbf{e}_{s,j+n} - \mathbf{e}_{s,j}\|_2 < \varepsilon \text{ or } \|\mathbf{e}_{s,j+n} - \mathbf{e}_{s,j}\|_2^2 < \varepsilon, \quad (\text{A.15})$$

when $\text{Lip} < 1$, which guarantees the existence of $\mathcal{P}_{\text{sp/unsp}}^{-1} = \mathcal{P}_{\text{unsp/sp}}$. It is worth mentioning that $\{\mathbf{e}_{s,j}\}$ is a Cauchy sequence [37] and (A.15) holds for $\text{Lip} < 1$ because of the Banach fixed-point theorem [38]. Furthermore, $\text{Lip} < 1$ is a weak condition/requirement to satisfy when choosing $\mathcal{P}_p$.

Hence, the proof of Theorem 1 is completed.

REFERENCES

- [1] H. Chen, B. Jiang, S. X. Ding, and B. Huang, "Data-driven fault diagnosis for traction systems in high-speed trains: A survey, challenges, and perspectives," *IEEE Trans. Intell. Transp. Syst.*, 2021, doi:10.1109/TITS.2020.3029946.
- [2] S. X. Ding, *Model-based Fault Diagnosis Techniques: Design Schemes, Algorithms, and Tools*, Berlin, Germany: Springer-Verlag, 2008.
- [3] D. Zhou, Y. Zhao, Z. Wang, X. He, and M. Gao, "Review on diagnosis techniques for intermittent faults in dynamic systems," *IEEE Trans. Ind. Electron.*, vol. 67, no. 3, pp. 2337–2347, Mar. 2020.
- [4] S. J. Qin, "Survey on data-driven industrial process monitoring and diagnosis," *Ann. Rev. Control*, vol. 36, no. 2, pp. 220–234, Dec. 2012.
- [5] B. Huang and R. Kadali, *Dynamic Modelling, Predictive Control and Performance Monitoring: A Data-Driven Subspace Approach*, London, U.K.: Springer-Verlag, 2008.
- [6] R. Isermann, "Model-based fault-detection and diagnosis—status and applications," *Ann. Rev. Control*, vol. 29, no. 1, pp. 71–85, Jan. 2005.
- [7] V. Venkatasubramanian, R. Rengaswamy, S. N. Kavuri, and K. Yin, "A review of process fault detection and diagnosis: Part III: Process history based methods," *Comput. Chem. Eng.*, vol. 27, no. 3, pp. 327–346, Mar. 2003.
- [8] S. Yin, S. X. Ding, X. Xie, and H. Luo, "A review on basic data-driven approaches for industrial process monitoring," *IEEE Trans. Ind. Electron.*, vol. 61, no. 11, pp. 6418–6428, Nov. 2014.
- [9] C. Dai, Z. Liu, K. Hu, and K. Huang, "Fault diagnosis approach of traction transformers in high-speed railway combining kernel principal component analysis with random forest," *IET Elect. Syst. Transp.*, vol. 6, no. 3, pp. 202–206, Sep. 2016.
- [10] H. Chen, B. Jiang, N. Lu, and W. Chen, *Data-driven Detection and Diagnosis of Faults in Traction Systems of High-speed Trains*, Cham, Switzerland: Springer Nature, 2020.
- [11] R. J. Patton and J. Chen, "Review of parity space approaches to fault diagnosis for aerospace systems," *J. Guid. Control Dyn.*, vol. 17, no. 2, pp. 278–285, Mar. 1994.
- [12] Z. Zhao, W. Jiang, and W. Gao, "Health evaluation and fault diagnosis of medical imaging equipment based on neural network algorithm," *Comput. Intell. Neurosci.*, 2021, doi:10.1155/2021/6092461.
- [13] S. X. Ding, *Advanced Methods for Fault Diagnosis and Fault-tolerant Control*, Berlin, Germany: Springer Nature, 2020.
- [14] J. Chen and R. J. Patton, *Robust Model-Based Fault Diagnosis for Dynamic Systems*, 1998, Norwell, MA, USA: Kluwer.
- [15] H. Wang, Z. Liu, A. Núñez, and R. Dollevoet, "Entropy-based local irregularity detection for high-speed railway catenaries with frequent inspections," *IEEE Trans. Instrum. Meas.*, vol. 68, no. 10, pp. 3536–3547, Oct. 2019.
- [16] S. X. Ding, *Data-driven Design of Fault Diagnosis and Fault-tolerant Control Systems*, London, U.K.: Springer-Verlag, 2014.
- [17] T. de Bruin, K. Verbert, and R. Babuška, "Railway track circuit fault diagnosis using recurrent neural networks," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 28, no. 3, pp. 523–533, Mar. 2017.
- [18] J. Seshadrinath, B. Singh, and B. K. Panigrahi, "Incipient interturn fault diagnosis in induction machines using an analytic wavelet-based optimized Bayesian inference," *IEEE Trans. Neural Netw. Learn. Syst.*, vol. 25, no. 5, pp. 990–1001, May. 2014.
- [19] Z. Liu, L. Wang, C. Li, and Z. Han, "A high-precision loose strands diagnosis approach for isoelectric line in high-speed railway," *IEEE Trans. Ind. Informat.*, vol. 14, no. 3, pp. 1067–1077, Mar. 2018.
- [20] T. Li, Z. Zhao, C. Sun, R. Yan, and X. Chen, "Multi-receptive field graph convolutional networks for machine fault diagnosis," *IEEE Trans. Ind. Electron.*, vol. 68, no. 12, pp. 12 739–12 749, Dec. 2021.
- [21] H. Chen, Z. Chai, O. Dogru, B. Jiang, and B. Huang, "Data-driven designs of fault detection systems via neural networks-aided learning," *IEEE Trans. Neural Netw. Learn. Syst.*, 2021, doi:10.1109/TNNLS.2021.3071292.
- [22] A. Bellet, A. Habrard, and M. Sebban, "Metric learning," *Synth. Lect. Artif. Intell. Mach. Learn.*, vol. 9, no. 1, pp. 1–151, Jan. 2015.
- [23] H. Chen, Z. Chen, Z. Chai, B. Jiang, and B. Huang, "A single-side neural network-aided canonical correlation analysis with applications to fault diagnosis," *IEEE Trans. Cybern.*, 2021, doi:10.1109/TCYB.2021.3060766.
- [24] Y. Yang, S. X. Ding, and L. Li, "Parameterization of nonlinear observer-based fault detection systems," *IEEE Trans. Autom. Control*, vol. 61, no. 11, pp. 3687–3692, Nov. 2016.
- [25] H. Chen, L. Li, C. Shang, and B. Huang, "Fault detection for nonlinear dynamic systems with consideration of modeling errors: A data-driven approach," *IEEE Trans. Cybern.*, 2021, to be published.

- [26] I. Mezić, “Koopman operator, geometry, and learning,” *arXiv*, Oct. 2020, arXiv:2010.05377.
- [27] A. L. Bruce, V. M. Zeidan, and D. S. Bernstein, “What is the Koopman operator? A simplified treatment for discrete-time systems,” in *Amer. Contr. Conf.*, Jul. 2019, pp. 1912–1917.
- [28] Y. Maki and K. A. Loparo, “A neural-network approach to fault detection and diagnosis in industrial processes,” *IEEE Trans. Control Syst. Technol.*, vol. 5, no. 6, pp. 529–541, Nov. 1997.
- [29] I. Rivals and L. Personnaz, “Black-box modeling with state-space neural networks,” *Neural Adapt. Control Technol.*, vol. 15, pp. 237–264, Apr. 1996.
- [30] S. Haykin, *Neural Networks and Learning Machines*, Hamilton, Canada: Pearson Education, 2010.
- [31] H. Chen, Z. Chai, B. Jiang, and B. Huang, “Data-driven fault detection for dynamic systems with performance degradation: A unified transfer learning framework,” *IEEE Trans. Instrum. Meas.*, 2021, doi:10.1109/TIM.2020.3033943.
- [32] A. N. Gomez, M. Ren, R. Urtasun, and R. B. Grosse, “The reversible residual network: Backpropagation without storing activations,” in *Proc. 31st Int. Conf. Neural Inf. Process. Syst.*, Dec. 2017, pp. 2211–2221.
- [33] J. Wang and S. J. Qin, “A new subspace identification approach based on principal component analysis,” *J. Process Control*, vol. 12, no. 8, pp. 841–855, Dec. 2002.
- [34] L. Ardizzone, J. Kruse, S. Wirkert, D. Rahner, E. W. Pellegrini, R. S. Klessen, L. Maier-Hein, C. Rother, and U. Köhe, “Analyzing inverse problems with invertible neural networks,” *arXiv preprint*, Feb. 2019, 1808.04730v3.
- [35] E. Sontag and Y. Wang, “Lyapunov characterizations of input to output stability,” *SIAM. J. Control Optim.*, vol. 39, no. 1, pp. 226–249, Jan. 2000.
- [36] K. I. Funahashi, “On the approximate realization of continuous mappings by neural networks,” *Neural Netw.*, vol. 2, no. 3, pp. 183–192, Mar. 1989.
- [37] S. Lang, *Algebraic Number Theory*, New York, USA: Springer Science & Business Media, 2013.
- [38] K. Ciesielski, “On Stefan Banach and some of his results,” *Banach J. Math. Anal.*, vol. 1, no. 1, pp. 1–10, Apr. 2007.



Hongtian Chen (Member, IEEE) received the B.S. and M.S. degrees in School of Electrical and Automation Engineering from Nanjing Normal University, China, in 2012 and 2015, respectively; and he received the Ph.D. degree in College of Automation Engineering from Nanjing University of Aeronautics and Astronautics, China, in 2019.

He had ever been a Visiting Scholar at the Institute for Automatic Control and Complex Systems, University of Duisburg-Essen, Germany, in 2018. Now he is a Post-Doctoral Fellow with the Department

of Chemical and Materials Engineering, University of Alberta, Canada. His research interests include process monitoring and fault diagnosis, data mining and analytics, machine learning, and quantum computation; and their applications in high-speed trains, new energy systems, and industrial processes.

Dr. Chen was a recipient of the Grand Prize of Innovation Award of Ministry of Industry and Information Technology of the People's Republic of China in 2019, the Excellent Ph.D. Thesis Award of Jiangsu Province in 2020, and the Excellent Doctoral Dissertation Award from Chinese Association of Automation (CAA) in 2020.



Zhigang Liu (Senior Member, IEEE) received the Ph.D. degree in Power system and its Automation from Southwest Jiaotong University, China in 2003. He is currently a Full Professor of the School of Electrical Engineering, Southwest Jiaotong University, Chengdu. He is also a Guest Professor of Tongji University, Shanghai.

He has authored three books and published more than 200 peer-reviewed journal and conference articles. His research interests include the electrical relationship of EMUs and traction, detection, and assessment of pantograph-catenary in high-speed railway. Dr. Liu is an Associate Editor-in-Chief of IEEE Transactions on Instrumentation and Measurement, Associate Editor of IEEE Transactions on Neural Networks and Learning Systems, and IEEE Transactions on Vehicular Technology. He received the IEEE TIM's Outstanding Associate Editors for 2019 and 2020, and the Outstanding Reviewer of IEEE Transactions on Instrumentation and Measurement in 2018.



Cesare Alippi (Fellow, IEEE) received the master's degree (cum laude) in electronic engineering and the Ph.D. degree from the Politecnico di Milano, Italy, in 1990 and 1995, respectively. He has been a Visiting Researcher with University College London, London, U.K.; the Massachusetts Institute of Technology, Cambridge, MA, USA; ESPCI, Paris, France; CASIA; and A*STAR, Singapore.

He is currently a Full Professor of information processing systems with the Politecnico di Milano and a Full Professor of cyber-physical and embedded

systems with the Università della Svizzera italiana, Lugano, Switzerland. He holds five patents. He has published in 2014 a monograph on Intelligence for Embedded Systems (Springer) and (co)authored more than 200 papers in international journals and conference proceedings. His current research activity addresses adaptation and learning in nonstationary environments and intelligence for embedded systems.

Dr. Alippi is a member and the chair of other IEEE committees. He received the IEEE Instrumentation and Measurement Society Young Engineer Award in 2004, the IBM Faculty Award in 2013, the 2016 IEEE TNNLS Outstanding Paper Award, and the 2016 International Neural Networks Society (INNS) Gabor Award. Among the others, he was the General Chair of the International Joint Conference on Neural Networks (IJCNN) in 2012 and the Program Chair and Co-Chair in 2014 and 2011, respectively. He was the General Chair of the IEEE Symposium Series on Computational Intelligence in 2014. He was an Associate Editor of the IEEE TRANSACTIONS ON INSTRUMENTATION AND MEASUREMENT and IEEE TRANSACTIONS ON NEURAL NETWORKS AND LEARNING SYSTEMS. He is a Distinguished Lecturer of the IEEE Computational Intelligence Society (CIS), a member of the Board of Governors of INNS, the Vice-President of IEEE CIS, an Associate Editor of the IEEE Computational Intelligence Magazine.



Biao Huang (Fellow, IEEE) received the B.Sc. and M.Sc. degrees in automatic control from the Beijing University of Aeronautics and Astronautics, Beijing, China, in 1983 and 1986, respectively, and the Ph.D. degree in process control from the University of Alberta, Edmonton, AB, Canada, in 1997.

In 1997, he joined the University of Alberta as an Assistant Professor with the Department of Chemical and Materials Engineering and is currently a Professor. He is the Natural Sciences and Engineering Research Council of Canada's Industrial Research Chair in Control of Oil Sands Processes and was also the Alberta Innovates Technology Futures Industry Chair in Process Control. He has applied his expertise extensively in industrial practice, particularly in the oil sands industry. His research interests include process control, system identification, control performance assessment, Bayesian methods, and state estimation.

Dr. Huang is a Fellow of the Canadian Academy of Engineering and the Chemical Institute of Canada. He was a recipient of Germany's Alexander von Humboldt Research Fellowship, the Canadian Chemical Engineer Society's Syncrude Canada Innovation and D. G. Fisher Awards, APEGA Summit Research Excellence Award, University of Alberta McCalla and Killam Professorship Awards, Petro-Canada Young Innovator Award, and a Best Paper Award from the Journal of Process Control.



Derong Liu (Fellow, IEEE) received the PhD degree in electrical engineering from the University of Notre Dame, USA, in 1994.

He became a Full Professor of Electrical and Computer Engineering and of Computer Science at the University of Illinois at Chicago in 2006. He was selected for the "100 Talents Program" by the Chinese Academy of Sciences in 2008, and he served as the Associate Director of The State Key Laboratory of Management and Control for Complex Systems at the Institute of Automation, from 2010

to 2016. He has published 13 books.

He received the International Neural Network Society's Gabor Award in 2018 and the IEEE CIS Neural Network Pioneer Award in 2022. He has been named a highly cited researcher consecutively for five years from 2017 to 2021 by Clarivate. He was the Editor-in-Chief of IEEE TRANSACTIONS ON NEURAL NETWORKS AND LEARNING SYSTEMS from 2010 to 2015. He is the Editor-in-Chief of Artificial Intelligence Review (Springer). He is a Fellow of the IEEE, a Fellow of the International Neural Network Society, a Fellow of the International Association of Pattern Recognition, and a Member of Academia Europaea (The Academy of Europe).