



A sparse dissimilarity analysis algorithm for incipient fault isolation with no priori fault information



Chunhui Zhao^{a,*}, Furong Gao^b

^a State Key Laboratory of Industrial Control Technology, College of Control Science and Engineering, Zhejiang University, Hangzhou, 310027, China

^b Department of Chemical and Biomolecular Engineering, The Hong Kong University of Science and Technology, Clear Water Bay, Kowloon, Hong Kong Special Administrative Region

ARTICLE INFO

Keywords:

Sparse dissimilarity
Distribution structure
Dissimilarity decomposition
Fault diagnosis
LASSO

ABSTRACT

The conventional multivariate statistical process control (MSPC) methods in general quantify the distance between the new sample and the modelling samples for fault detection and diagnosis, which, however, do not check the changes of data distribution as long as monitoring statistics stay inside normal region enclosed by control limit and thus are not sensitive to incipient changes. In the present work, a sparse dissimilarity (SDISSIM) algorithm is developed which can isolate the incipient abnormal variables that change the data distribution structure and does not need any priori fault knowledge. First, the distribution dissimilarity is decomposed deeply and significant dissimilarity is extracted to integrate the critical difference of variable covariance structure between the reference normal operation distribution and the actual distribution. Second, a sparse regression-based optimization problem is formulated to isolate abnormal variables associated with changes of distribution structure. Sparse coefficients are obtained with only a small fraction of variables' coefficients nonzeros, pointing to abnormal variables. As illustrations, SDISSIM is applied to both simulated and real industrial process data with encouraging results to figure out the slight distortions.

© 2017 Elsevier Ltd. All rights reserved.

1. Introduction

Modern industrial processes have shown an urgent demand for keeping process safety and improving product quality. Efficient process monitoring has been an important issue and is drawing increasing attentions with studies of fault detection and diagnosis (Cauffriez, Grondel, & Loslever, 2016; Chiang, Russell, & Braatz, 2000; Kariwala, Odiowei, Cao, & Chen, 2010; Kerkhof, Vanlaer, Gins, & Impe, 2013; Portnoy, Melendez, & Pinzon, 2016; Tong, El-Farra, Palazoglu, & Yan, 2014; Undey & Cinar, 2002; Yu & Rashid, 2013; Zhao, Sun, & Gao, 2012; Zhao, Yao, Gao, & Wang, 2010; Zhang, Zhao, Wang, & Wang, 2017). Among them, principal component analysis (PCA) (Jackson, 2005; Wold, Esbensen, & Geladi, 1987) and partial least squares (PLS) (Burnham, Viveros, & MacGregor, 1996; Dayal & MacGregor, 1997) as the typical representation of multivariate statistical process control (MSPC) methods have made an impact since the last three decades because their derivation requires a minimal *a priori* knowledge about process physics. They in general use distance-based statistics and the job is to timely detect any deviation from normal or “in-control” region that has been defined to accommodate the acceptable variations. It is

extremely important to detect the early occurrence of a small process deviation before the serious failure of the overall process so that it can be handled to prevent undesired consequences in which fault isolation is a main function.

Fault detection and isolation are two essential steps as stated by Chiang, Kotanchek, and Kordon (2004). For fault detection and isolation task, it requires that information about the normal (fault-free) behaviour should be available and the main principle behind is to compare system's actual behaviour against its nominal one to check whether the consistency stays and what distorts the consistency. The fault isolation follows the fault detection task to identify the variables primarily responsible for the disturbance, making the subsequent step of root cause diagnosis easier. Data-driven fault diagnosis methods are reviewed in Isermann (1997, 2006); MacGregor and Cinar (2012); Qin (2012); Russell and Braatz (2001); Venkatasubramanian, Rengaswamy, Kavuri, and Yin (2003); and Yin, Ding, Xie, and Luo (2014); and various methods are compared. In general, the fault isolation methods can be classified into two typical classes. One class is the approach that does not

* Corresponding author. Fax: +86 571 87951879.
E-mail address: chzhao@zju.edu.cn (C. Zhao).

need any historical fault information. The typical representation is contribution plot (Alcala & Qin, 2009; Li, Alcala, Qin, & Zhou, 2011; Miller, Swanson, & Heckler, 1998; Westerhuis, Gurden, & Smilde, 2000) that is a standard tool to isolate faulty variables without a priori knowledge by determining the contribution of each variable to the fault detection statistics. Despite of its simplicity, however, the principle may not be correct because it assumes that the fault-free variables after fault occurs follows the same distribution as that under normal operating conditions. The other class is the approach that requires exploration of historical fault data for extraction of fault feature. The typical representation is discriminant analysis based classification method (Dunia & Qin, 1998) and reconstruction based fault diagnosis (He, Qin, & Wang, 2005; Zhao & Sun, 2013) method. Zhao and Sun (2013) proposed the idea of relative changes to improve the fault reconstruction model and thus remove out-of-control monitoring statistics more efficiently. Instead of directly modelling the fault data, the relative changes from normal to fault condition are analysed along monitoring directions derived from normal data so that the significant directions responsible for out-of-control monitoring statistics are extracted to reconstruct disturbances. Further, faulty variable selection strategy (Zhao & Gao, 2017; Zhao & Wang, 2016) has been developed in combination with reconstruction technique for online fault diagnosis, in which, the key faulty variables can be isolated to better describe fault characteristics and thus improve fault diagnosis performance. However, sufficient historical fault data have served as a basis for fault reconstruction and classification, which, however, may be difficult to be obtained in practice. Besides, both of them may have difficulty handling unknown disturbances that are not covered by historical fault data.

The above fault diagnosis methods in general calculate the distance index to evaluate the derivations from the normal condition. Therefore, they may not always function well and in particular they may not be sensitive to incipient changes in the variable covariance structure or changes in the geometry of the underlying distribution decomposition. Recently, some work have noticed this problem and been presented to address this issue. Kano, Hasebe, Hashimoto, and Ohno (2001a, b) calculated a new index to evaluate the changes in the direction of each principal component. However, it cannot detect changes in the process variance. Kruger, Kumar, and Littler (2007) addressed this issue by incorporating the local approach (Basseville, 1998) into the multivariate statistical monitoring framework and constructed two univariate statistics to detect changes in the directions of the eigenvectors that span the model plane and in the eigenvalues that represent process variance. However, Kruger et al. only simply plotted the difference of covariance between the process reference data and one fault condition and used departures of elements to indicate which variables were mostly affected by a fault condition. Kano, Hasebe, Hashimoto, and Ohno (2002) proposed a dissimilarity index, referred to as DISSIM method, to quantitatively evaluate the distribution difference between normal condition and fault condition. DISSIM method is based on the idea that a change of operating condition can be detected by monitoring the distribution of time-series data covering both distribution directions and variances, which reflects the corresponding operating condition. One pattern recognition algorithm, Karhunen–Loeve (KL) expansion (Fukunaga & Koontz, 1970), was introduced to make transformation so that the concerned two data sets shared the same eigenvectors, i.e., the same directions. Thus the distribution difference can be simply reflected by the difference between eigenvalues, i.e., the process variances. It has been generalized for fault detection of batch processes (Zhao, Wang, & Jia, 2007) and the nonlinear expansion of DISSIM method was also reported by Zhao, Wang, and Zhang (2009) for fault detection of nonlinear processes. However, although the DISSIM method has been successfully used for detection of incipient faults, fault isolation of the abnormal variables that distort the variable covariance structure has not been well addressed. Like the conventional contribution plot, a contribution of each process variable to the dissimilarity index was calculated by Kano et al. (2001a, b) to simply identifying the abnormal variables, which thus has the same problem as contribution plot.

This paper proposes a sparse dissimilarity (SDISSIM) algorithm for online incipient fault diagnosis without priori fault information. Considering the power of DISSIM algorithm for distribution monitoring, it is used as the basic analysis algebra from which the faulty variable isolation strategy is developed that incorporates the DISSIM approach into fault diagnosis framework. First, it gives rise to further decomposition of dissimilarity directions which can integrate the major changes in the data distribution structure that are mostly affected by an incipient fault. Second, a sparse regression-type optimization is formulated to obtain a compact set of potential faulty variables for the purpose of isolating disturbed variables that have caused changes in the variable covariance structure resulting from the presence of an incipient fault. The paper demonstrates that the new method is more sensitive to isolation of incipient faulty variables that are responsible for distortion of the underlying covariance structure.

The major contribution is summarized as below:

- (1) The proposed method is for incipient fault diagnosis from the perspective of changes of process distribution.
- (2) A sparse DISSIM algorithm is formulated which can automatically online isolate multiple incipient fault variables that cause changes of process distribution without any priori fault information.
- (3) A judgement strategy is developed to initially determine whether all possible faulty variables that are responsible for the concerned incipient fault have been online identified without having trouble recalculating the monitoring statistics.

2. DISSIM revisit and motivation

DISSIM method (Kano et al., 2002) can quantify the distribution difference between two data sets based on the index of dissimilarity analysis which incorporates a classification method based on Karhunen–Loeve (KL) expansion (Fukunaga & Koontz, 1970) into the MSPC framework. The analysis subjects are two data sets, which share the same number of variables but have different number of samples. One data set is referred to be the reference which is normalized to have zero-mean and unit-variance. And the normalization information can then be employed to deal with the other data set, which thus covers the distribution information departing from the reference one. The distribution difference between the two data sets is evaluated for modelling. The details of DISSIM algorithm can refer to the work by Kano et al. (2002). Here, their modelling procedure is simply outlined as follows.

- (1) Calculate the covariance matrix of the mixture of two data sets.
- (2) Perform eigenvalue decomposition on the covariance matrix.
- (3) Transform the original data matrices by the eigenvectors obtained from Step (2).
- (4) Calculate the covariance matrices of the transformed data matrices.
- (5) Conduct eigenvalue decomposition on each transformed covariance matrix.
- (6) Calculate the index D to quantify the distribution dissimilarity and determine the control limit.

For on-line fault detection, the current data matrix representing the actual operating condition is updated continuously by moving the time-window forward step-wise, and it is normalized using the mean and the variance information obtained from the reference data. Then, the dissimilarity index D is calculated to evaluate the distribution difference between the actual and the reference data sets. If the index is outside the control limit, the current operation condition is judged to present different covariance structure from the reference one.

By DISSIM algorithm, it first transforms the concerned two data sets resulting in the same eigenvectors, i.e., the same distribution directions.

Download English Version:

<https://daneshyari.com/en/article/5000261>

Download Persian Version:

<https://daneshyari.com/article/5000261>

[Daneshyari.com](https://daneshyari.com)