

A Robust Discriminate Analysis Method for Process Fault Diagnosis

D. Wang* and J. A. Romagnoli

Dept. of Chemical Engineering, the University of Sydney, NSW 2006, Australia

Abstract:

A robust Fisher discriminant analysis (FDA) strategy is proposed for process fault diagnosis. The performance of FDA based fault diagnosis procedures could deteriorate with the violation of the assumptions made during conventional FDA. The consequence is a reduction in accuracy of the model and efficiency of the method, with the increase of the rate of misclassification. In the proposed approach, an M-estimate winsorization method is applied to the transformed data set; this procedure eliminates the effects of outliers in the training data set, while retaining the multivariate structure of the data. The proposed approach increases the accuracy of the model when the training data is corrupted by anomalous outliers and improves the performance of the FDA based diagnosis by decreasing the misclassification rate. The performance of the proposed method is evaluated using a multipurpose chemical engineering pilot-facility.

Key Words: discriminant analysis, robustness, fault diagnosis, and process monitoring.

1. Introduction

Chemical processes experience abnormal conditions that may lead to out-of-specification products or even process shutdown. These abnormal conditions are often related to the same root causes. Data driven process fault diagnosis techniques are often employed in process industries due to their ease of implementation, requiring very little modelling effort and *a priori* information. Given that there are multiple datasets in the historical database, each associated with a different abnormal condition (root cause), the objective of fault diagnosis is to assign the on-line out-of-control observations to the most closely related fault class.

Fisher discriminate analysis (FDA) is a superior linear pattern classification technique, which has been applied in industry for fault diagnosis (Russell et. al. 2000). By maximising the scatter between classes and minimising the scatter within classes, FDA projects faulty data into a feature space so that data from different classes are maximally separated. Discriminant functions are established associated with the feature space so that the classification of new faulty data is undertaken by projecting it into the feature space and comparing their scores. As a dimensionality reduction technique for feature extraction as PCA, FDA is superior to PCA because it takes into account the information between the classes and is well suited for fault diagnosis. FDA also has better performance than other techniques such as KNN and SIMCA (Chiang, et. al. 2004).

* to whom correspondence should be addressed: dwang@chem.eng.usyd.edu.au

Even through the above advantages, there are still unsolved issues within the application of FDA approaches. One key aspect is the robustness of the approach when dealing with real data. It is known that, in FDA, the most difficult assumption to meet is the requirement for a normal distribution on the discriminating variables, which are formed by measurements at interval level. Practical examination tells us that the real plant data seldom satisfy to this crucial assumption. The data are usually unpredictable having, for example, heavier tails than the normal ones, especially when data contain anomalous outliers. This will inevitably result in the loss of performance leading in some cases to wrong modelling in the feature extraction step, which in turn leads to misclassifications of the faulty conditions.

In this paper, a robust discriminant analysis method for process fault is presented. In the proposed approach, without eliminating the data in the training set, robust estimations of within-scatter matrix and between-class-scatter matrix are obtained using reconstructed data using M-estimator theory. A winsorization process is applied in the score space, which eliminates the effects of outliers in the original data in the sense of maximum likelihood estimation. The robust estimator used in this work is based on the Generalised T distribution, which can adaptively transform the data to eliminate the effects of the outliers in the original data (Wang et. al. 2003, Wang et. al., 2004). Consequently, a more accurate model is obtained and this procedure is optimal in the sense of minimising the number of misclassifications for process fault diagnosis.

2. Process Fault Diagnosis Using Discriminant Analysis

2.1 Discriminant analysis

Let the training data for all faulty classes be stacked into a n by m matrix $X \in \mathbb{R}^{n \times m}$, where n is the observation number and m is the variable number. The within-class-scatter matrices S_w and the between-class-scatter matrix S_b contain all the basic information about the relationship within the groups and between them (Russell et. al. 2000). The FDA can be obtained by solving the generalized eigenvalue problem: $S_b u_k = \lambda_k S_w u_k$, where λ_k indicates the degree of overall separation among the classes by projecting the data onto new coordinate system represented by u_i . After the above process, FDA decomposes the observation $X \in \mathbb{R}^{n \times m}$ as

$$X = TU^T = \sum_{i=1}^m t_i u_i^T \quad (1)$$

2.2 Process fault diagnosis based on FDA

After projecting data onto the discriminant function subspace, the data of different groups will cluster around their centroids. The objective of fault diagnosis is to assign the on-line out-of-control observations to the most closely related fault classes using classification techniques. An intuitive means of classification is to measure the distances from the individual case to each of the group centroids and classify the case into the closest group. Considering the fact that, in the chemical engineering measurements there are correlated variables, different measurement units, and different standard deviations, the concept of distance needs to be well defined. A generalized distance measure is introduced (Mahalanobis distance): $D^2(x_i | G_k) = (x_i - \bar{x}_k)^T V_k^{-1} (x_i - \bar{x}_k)$, where

$D^2(x_i | G_k)$ is the squared distance from a specific case x_i to $\bar{x}_k$, the centroid of group k , where V_k is the sample covariance matrix of group k . After calculating D^2 for each group, one would classify the case into the group with the smallest D^2 , that group is the one in which the typical profile on the discriminating variables most closely resembles the profile of this case.

By classifying a case into the closest group according to D^2 , one is implicitly assigning it into the group for which it has the highest probability of belonging. If one assumes that every case must belong to one of the groups, one can compute a probability of group membership for each group: $P(G_k | x_i) = P(x_i | G_k) / \sum_{i=1}^g P(x_i | G_i)$. This is a posterior probability; the classification on the largest of these values is also equivalent to using the smallest distance.

3. Robust Discriminant Analysis Based on M-estimate Winsorization

The presence of outliers in the training data will result in deviations of discriminant function coefficients from the real ones, so that the coordinate system for data projection may be changed. Fault diagnosis based on this degraded model will inevitably increase the misclassification rate. A robust remedy procedure is proposed here, to reduce the effects of outliers in the training data. After implementing FDA, the outliers in the original data $X \in \mathbb{R}^{n \times m}$ can manifest themselves in the score space. By recurrently winsorizing the scores and replacing them with suitable values, it is possible to detect multivariate outliers and replace them by values which conform to the correlation structure in the data.

3.1 Winsorization

Consider the linear regression problem: $y = f(X, \theta) + \varepsilon$, where: y is a $n \times 1$ vector of dependent variables, X is a $n \times m$ matrix of independent variables, and θ is a $p \times 1$ vector of parameters, ε is a $n \times 1$ vector of model error or residual. An estimation of parameter θ ($\hat{\theta}$) can be obtained by optimization or least squares method. With the parameter $\hat{\theta}$ estimated, the prediction or estimation of the dependent variable $y_i (i = 1, \dots, n)$ is given by $\hat{y}_i = f(x_i, \hat{\theta})$ and the residual is given by $r_i = y_i - \hat{y}_i$. In the winsorization process, the variable y_i is transformed using pseudo observation according to specified M-estimates, which characterizes the residual distribution. The normal assumption of residual data will result in poor performance of winsorization. In this work, we will fit the residual data to a more flexible distribution, i.e. the generalized T distribution, which can accommodate the shapes of most distributions one meets in practice, and then winsorize the variable y_i using its corresponding influence function.

3.2 Robust discriminant analysis based on M-estimate winsorization

The proposed robust estimator for FDA modelling is based on the assumption that the data in the score space follow the generalized T distribution (GT) (Wang and Romagnoli, 2003), which has the flexibility to accommodate various distributional shapes:

$$f_{GT}(u; \sigma, p, q) = \frac{p}{2\sigma q^{1/p} B(1/p, q) \left(1 + |u|^p / q \sigma^p\right)^{q+1/p}} \quad -\infty < u < \infty \quad (2)$$

where: u is the argument, σ, p, q are distributional parameters, σ corresponds to the standard deviation, p and q are parameters corresponding to the shape of distribution. This density is symmetric about zero, uni-modal, and suitable to describe the error characteristics in most cases. By choosing different values of p and q , the GT distribution will accommodate the real shape of the error distribution. The tail behaviour and other characteristics of the distribution depend upon these two distributional parameters, which can be estimated from the data (Wang et. al., 2003). The robustness of the estimator, based on a GT distribution, can be discussed by investigating its influence function (ψ), which is given by

$$\psi_{GT}(u; \sigma, p, q) = (pq + 1) \text{sign}(u) |u|^{q-1} / (q\sigma^p + |u|^p) \quad (3)$$

The technique of winsorization will be used in FDA to eliminate the effects of outliers in the following way. The data value y in score space can be transformed into a new value y^w by winsorization. By doing this, the large values exhibited as outliers in the original data set will be brought closer to the other observations after they are transformed from the score space back to the original data space according Eq. (1). A new FDA model is obtained using the new data set. This process is carried out iteratively until there is not much change in the discriminant function coefficient.

4. Case study

The proposed robust FDA (RFDA) strategy is applied to a case study of a pilot-scale setting containing two CSTRs, a mixer and a number of heat exchangers (Wang et. al., 2004). A process flowsheet is presented in Figure 1. The total number of measurement variables in this study is 27, including 12 temperature variables, 11 flowrate variables and 4 level variables. Plant data are combined with simulations to generate data with different distributions and outliers. This architecture is used to generate data under normal and faulty conditions. To investigate the proficiency and performance of the proposed robust approach RFDA compared to the conventional FDA (CFDA), training data that consist of three faulty classes are generated when the plant is running under three different faulty conditions, and the data are contaminated with outliers as well as

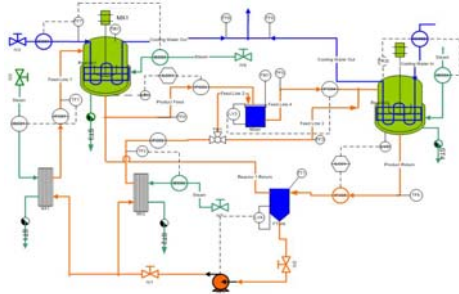


Figure 1. Two CSTR's Flowsheet

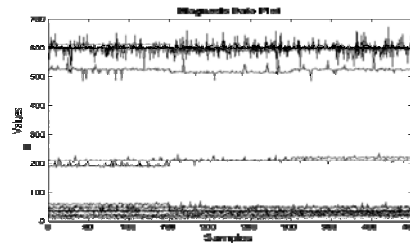


Figure 2. Corrupted Training Data (Class 1: Sample 1-150; Class 2: Sample 151-300; Class 3: Sample 301-450)

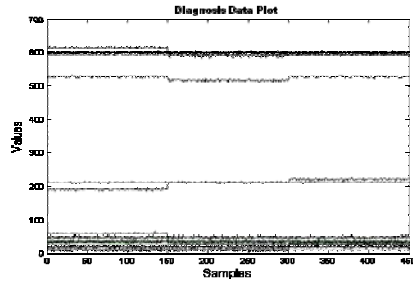


Figure 3. Uncorrupted Training Data (Class 1: Sample 1-150; Class 2: Sample 151-300; Class 3: Sample 301-450)

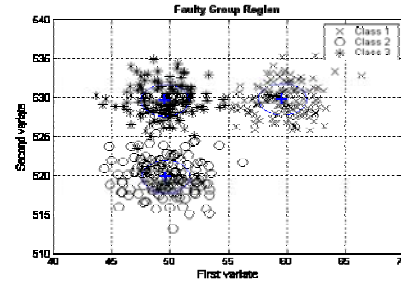


Figure 4. The projections of training data (without outliers) onto the first two loading vectors using CFDA.

normal noise (10% noise) as shown in Figure 2. The training data without outliers are also plotted for comparison (Figure 3).

For this corrupted training data set in Figure 2, both CFDA and RFDA are applied to build the model. The calculation results show that different FDA models are obtained by these two different approaches. For the same number of discriminant function retained (for example, 2 in this study), the different discriminating powers of the functions are presented in the two approaches, with the ones in RFDA having larger discriminating power than the CFDA one. This can be illustrated by investigating the eigenvalues from these two approaches tabulated in Table 1. It is custom in FDA that the nonzero eigenvalues are presented in the order of descending magnitude, and the size of the eigenvalue is related to the discriminating power of that function: the larger the eigenvalue, the greater the discrimination. From the table one can imagine that RFDA gives better separation of the faulty groups than CFDA, because the eigenvalues from RFDA are much larger than that from CFDA.

By projecting the training data to the reduced discriminant function subspace, one can easily visualize the separation degrees of the approaches. Figure 4 shows the projections of training data (without outliers) for three faulty classes onto the first two loading vectors using CFDA. One can see that there are three clusters of data, which are representing three different faulty groups. The centroids are represented by the cross signs inside the circles. However, when the training data are contaminated with outliers (data in Figure 2), the projected data using CFDA will widely spread in the reduced subspace and there are great overlapping between the group regions (Figure 5). This indicates that the outliers in the training set have a great influence on the FDA model and the misclassification rate based on the less accurate model will inevitably increase. Using RFDA, to the corrupted data set, will reduce the effects of the outliers and makes the model close to the real one. Figure 6 illustrates the projection of training data using RFDA. From the figure, one can see that faulty groups are well separated and the locations and the ranges of clusters are nearly the same as the case when there are no contaminations and CFDA is used (comparing Figure 4). The rationale behind this is that the robust winsorization process in the discriminant functions subspace brings the spread data, those which manifest outliers in the original data space, to the centroids of the clusters, so that this will eliminate the effects of outliers and better model is expected. The misclassification rates of these two approaches are reported in Table 2 when they are used for fault diagnosis using 150 on-line fresh data in each group. From

the table one can see that diagnosis, using CFDA gives a large misclassification rate, while diagnosis based on RFDA has similar misclassification rate with the situation where the model is built by CFDA using uncorrupted data.

5. Conclusions

A robust FDA modelling approach (RFDA) based on winsorization in score space using adaptive robust estimator was developed and presented. The effects of outliers in the training data can be eliminated by the method while the effectiveness as well as the robustness is retained by using GT-like estimator. It was shown that the proposed method can obtain a more accurate model when the normal assumption is violated in the training data. Process fault diagnosis by RFDA has better performance than that of the conventional FDA, resulting less misclassification rate.

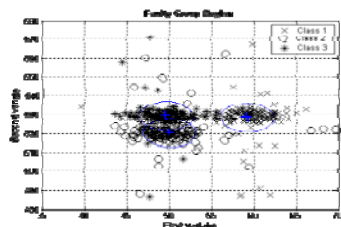


Figure 5. The projections of training data (with outliers) onto the first two loading vectors using CFDA.

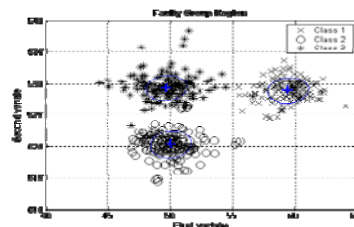


Figure 6. The projections of training data (with outliers) onto the first two loading vectors using RFDA.

Table 1 Discriminant Ability (λ_i : i^{th} eigenvalue)

Eigenvalues	RFDA	CFDA
λ_1	0.8401	0.1761
λ_2	0.2918	0.0678

Table 2 Misclassification Rates of the Approaches for Testing Data

Misclassification Rate	RFDA(with outliers)	CFDA(with outliers)	CFDA(without outliers)
Class 1	0.040	0.240	0.033
Class 2	0.033	0.340	0.040
Class 3	0.020	0.253	0.046
Overall	0.031	0.277	0.040

References

- Chiang, L. H., M. E. Kotanchek and A. K. Kordon (2004), "Fault diagnosis based on Fisher discriminate analysis and support vector machines", *Comp. & Chem. Eng.*, **28**, 8, 1389-1401
- Devlin S. J., Guanadesikan, R. and Kettenring, J. R. (1981), "Robust Estimation of Dispersion matrices and principal components", *J. of Amer. Stats. Asso.* 76, 354-362
- Hoo, K.A., K.J. Tvarlapati, M.J. Piovoso and R. Hajare (2002), "A Method of Robust Multivariate Outliers Replacement", *Comp. Chem. Eng.*, 26, 17-39
- Russell, Evan L., L. H. Chiang and R. D. Braatz (2000), "Data-driven Techniques for Fault Detection and Diagnosis in Chemical Processes", Springer.UK
- Wang, D. and J. A. Romagnoli (2003), "A Framework of Robust Data Reconciliation Based on a Generalized Objective Function", *Ind. Eng. Chem. Res.*, Vol. 42, No. 13, pp. 3075-3084.
- Wang, D. and J. A. Romagnoli (2004), "A Robust PCA Modelling Method for Process Monitoring", 7th International Symposium on Advanced Control of Chemical Processes (ADCHEM), 11-14 Jan. 2004 Hong Kong, China, 417-422.
- Wang, D. and J. A. Romagnoli (2004), "Robust Multi-Scale Principal Components Analysis with Applications to Process Monitoring", *J. of Process Control*. In process