



# Détection et Localisation de Défauts par Analyse en Composantes Principales

## THÈSE

présentée et soutenue publiquement le 30 Juin 2003

pour l'obtention du

**Doctorat de l'Institut National Polytechnique de Lorraine**  
(spécialité automatique et traitement numérique du signal)

par

**HARKAT Mohamed-Faouzi**

### Composition du jury

|                      |                                                   |                                                                                                                                                                                                 |
|----------------------|---------------------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| <i>Président :</i>   | M. MONSION                                        | Professeur à l'Université de Bordeaux (LAP)                                                                                                                                                     |
| <i>Rapporteurs :</i> | G. DELAUNAY<br>T. DENOEUX                         | Professeur à l'Université de Reims<br>Professeur à l'Université de Compiègne                                                                                                                    |
| <i>Examineurs :</i>  | A. RICHARD<br>J. RAGOT<br>G. MOUROT<br>D. PEARSON | Professeur à l'Université Henri Poincaré, Nancy<br>Professeur à l'INPL, Nancy (CRAN)<br>Ingénieur de recherche à l'INPL, Nancy (CRAN)<br>Professeur à l'IUT de Roanne, Université de St-Etienne |

Mis en page avec la classe thloria.

*À mes chers parents.*

*À ma femme et ma petite meriem.*



## Remerciements

Les travaux présentés dans ce mémoire ont été réalisés au sein de l'équipe Diagnostic et Robustesse du Centre de Recherche en automatique de Nancy (UMR-CNRS 7039).

Je tiens à remercier Monsieur le Professeur José RAGOT pour sa constante disponibilité et ses précieux conseils qui ont permis à ce travail de voir le jour.

Mes remerciements vont à Gilles MOUROT pour son aide, sa disponibilité, ses judicieux conseils pendant toute la durée de ma thèse. Ses remarques m'ont été d'une grande utilité dans l'avancement des travaux.

Je remercie Messieurs les Professeurs Thierry DENOEU et Georges DELAUNAY d'avoir accepté de rapporter mon mémoire et pour l'intérêt qu'ils ont bien voulu porter à ce travail. Leur lecture approfondie du mémoire, leurs remarques et interrogations judicieuses m'ont été très précieuses.

Mes remerciements s'adressent également à Monsieur le Professeur Michel MONSION d'avoir accepté d'examiner ce travail et d'assurer la présidence du jury. Mes remerciements vont à Messieurs les Professeurs David William PEARSON et Alain RICHARD qui ont accepté d'examiner et de juger ce travail.

Je n'oublie pas dans mes remerciements tous les membres de l'équipe Diagnostic et Robustesse du Centre de recherche en Automatique de Nancy pour la bonne ambiance qu'ils ont su faire régner au sein de l'équipe et tout particulièrement Marjorie Schwartz pour sa constante disponibilité.

Enfin, je ne saurais oublier de trop remercier mes parents pour leur soutien le long de ce parcours.



# Table des matières

|                                                                           |             |
|---------------------------------------------------------------------------|-------------|
| <b>Table des figures</b>                                                  | <b>ix</b>   |
| <b>Notations</b>                                                          | <b>xiii</b> |
| <b>Chapitre 1 Introduction générale</b>                                   | <b>1</b>    |
| <b>Chapitre 2 Analyse en Composantes Principales Linéaires</b>            | <b>7</b>    |
| 2.1 Introduction . . . . .                                                | 7           |
| 2.2 Principes de l'analyse en composantes principales . . . . .           | 9           |
| 2.3 Identification du modèle ACP . . . . .                                | 14          |
| 2.3.1 Estimation des paramètres du modèle . . . . .                       | 14          |
| 2.3.2 Détermination de la structure du modèle . . . . .                   | 18          |
| 2.3.2.1 Critères Heuristiques . . . . .                                   | 18          |
| 2.3.2.2 Critères de généralisation . . . . .                              | 19          |
| 2.3.2.3 Critère de validation croisée . . . . .                           | 19          |
| 2.3.2.4 Nombre de composantes pour une meilleure reconstruction . . . . . | 20          |
| 2.3.3 Sélection de capteurs . . . . .                                     | 25          |
| 2.4 Conclusion . . . . .                                                  | 28          |
| <b>Chapitre 3 Détection et localisation de défauts par ACP</b>            | <b>33</b>   |
| 3.1 Introduction . . . . .                                                | 33          |
| 3.2 Détection de défauts . . . . .                                        | 35          |
| 3.2.1 Génération de résidus par estimation d'état . . . . .               | 35          |
| 3.2.1.1 Statistique $SPE$ . . . . .                                       | 40          |
| 3.2.1.2 Statistique $T^2$ de Hotelling . . . . .                          | 42          |
| 3.2.1.3 Statistique SWE . . . . .                                         | 43          |
| 3.2.1.4 Statistique combinée . . . . .                                    | 44          |
| 3.2.1.5 Filtrage EWMA pour la détection . . . . .                         | 45          |
| 3.2.1.6 Statistique $D$ proposée . . . . .                                | 49          |
| 3.2.2 Génération de résidus par estimation paramétrique . . . . .         | 56          |
| 3.3 Localisation de défauts . . . . .                                     | 59          |

|                   |                                                                                              |            |
|-------------------|----------------------------------------------------------------------------------------------|------------|
| 3.3.1             | Localisation par structuration des résidus . . . . .                                         | 60         |
| 3.3.1.1           | Structuration de résidus . . . . .                                                           | 61         |
| 3.3.1.2           | Structuration des résidus avec l'ACP . . . . .                                               | 62         |
| 3.3.1.3           | Résidus structurés avec maximisation de la sensibilité aux défauts . . . . .                 | 63         |
| 3.3.2             | Localisation utilisant un banc de modèles . . . . .                                          | 69         |
| 3.3.2.1           | Localisation par ACP partielles . . . . .                                                    | 70         |
| 3.3.2.2           | Localisation par élimination . . . . .                                                       | 72         |
| 3.3.2.3           | Localisation basée sur le principe de reconstruction . . . . .                               | 74         |
| 3.3.2.4           | Conditions nécessaires et suffisantes pour la localisation par re-<br>construction . . . . . | 82         |
| 3.3.2.5           | Localisation par reconstruction utilisant l'indice $D_i$ . . . . .                           | 83         |
| 3.3.3             | Localisation par calcul des contributions . . . . .                                          | 85         |
| 3.3.4             | Localisation de défauts multiples . . . . .                                                  | 92         |
| 3.4               | Identification des caractéristiques du défaut . . . . .                                      | 97         |
| 3.5               | Conclusion . . . . .                                                                         | 98         |
| <b>Chapitre 4</b> | <b>Analyse en composantes principales non linéaires (ACPNL)</b>                              | <b>101</b> |
| 4.1               | Introduction . . . . .                                                                       | 101        |
| 4.2               | Analyse en composantes principales non linéaires . . . . .                                   | 103        |
| 4.3               | Méthode des Courbes Principales . . . . .                                                    | 104        |
| 4.3.1             | Algorithme de calcul des courbes principales de Hastie . . . . .                             | 107        |
| 4.3.1.1           | Etape de projection . . . . .                                                                | 107        |
| 4.3.1.2           | Etape d'estimation . . . . .                                                                 | 107        |
| 4.3.2             | Algorithme de calcul des courbes principales de Verbeek . . . . .                            | 107        |
| 4.4               | Approches neuronales de l'ACPNL . . . . .                                                    | 110        |
| 4.4.1             | ACPNL par réseau à cinq couches . . . . .                                                    | 111        |
| 4.4.1.1           | Apprentissage du réseau . . . . .                                                            | 112        |
| 4.4.2             | ACPNL par optimisation des entrées du réseau (Input Training network) . . . . .              | 116        |
| 4.4.3             | Réseaux à trois couches et courbes principales . . . . .                                     | 117        |
| 4.4.4             | Réseaux de fonctions à base radiale (RBF) . . . . .                                          | 119        |
| 4.4.4.1           | Réseau RBF avec maximisation de la variance de $t$ . . . . .                                 | 120        |
| 4.4.4.2           | Approche combinant les courbes principales et les réseaux RBF . . . . .                      | 123        |
| 4.4.4.3           | Détermination du nombre de composantes non linéaires . . . . .                               | 124        |
| 4.5               | Conclusion . . . . .                                                                         | 129        |
| <b>Chapitre 5</b> | <b>Détection et localisation de défauts capteurs d'un RSQL</b>                               | <b>131</b> |
| 5.1               | Introduction . . . . .                                                                       | 131        |
| 5.2               | Description du phénomène . . . . .                                                           | 132        |

---

|                                               |                                                                        |            |
|-----------------------------------------------|------------------------------------------------------------------------|------------|
| 5.3                                           | Description du réseau de surveillance de la qualité de l'air . . . . . | 133        |
| 5.4                                           | Prétraitement des données . . . . .                                    | 136        |
| 5.5                                           | Application de l'ACP linéaire . . . . .                                | 139        |
| 5.6                                           | Conclusion . . . . .                                                   | 151        |
| <b>Chapitre 6 Conclusions et Perspectives</b> |                                                                        | <b>159</b> |
| 6.1                                           | Conclusion . . . . .                                                   | 159        |
| 6.2                                           | Perspectives . . . . .                                                 | 161        |
| <b>Références bibliographiques</b>            |                                                                        | <b>163</b> |
| <b>Bibliographie</b>                          |                                                                        | <b>165</b> |
|                                               |                                                                        | <b>173</b> |



## Table des figures

|      |                                                                                                                                                                                                                                                    |    |
|------|----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 1.1  | Schéma de principe de la détection et de la localisation de défauts à base de modèles analytiques . . . . .                                                                                                                                        | 2  |
| 1.2  | Exemple de résidu directionnel ou le résidu est dans la direction du défaut 1 . .                                                                                                                                                                  | 3  |
| 1.3  | Principe de génération des résidus . . . . .                                                                                                                                                                                                       | 4  |
| 1.4  | Structure générale d'un système de détection . . . . .                                                                                                                                                                                             | 4  |
| 2.1  | Déroulement d'une analyse en composantes principales. (a) Distribution d'entrée. (b) Centrage et réduction de cette distribution. (c) Les deux axes principaux, correspondant aux vecteurs propres de la matrice de covariance de la distribution. | 13 |
| 2.2  | Mesures simulées de $x_1$ , $x_4$ et $x_7$ du premier exemple d'illustration . . . . .                                                                                                                                                             | 16 |
| 2.3  | Evolution des composantes principales de l'exemple 1 . . . . .                                                                                                                                                                                     | 17 |
| 2.4  | Principe de la méthode itérative de reconstruction . . . . .                                                                                                                                                                                       | 20 |
| 2.5  | Evolution des variables $x_1$ , $x_5$ et $x_9$ de l'exemple 2 . . . . .                                                                                                                                                                            | 25 |
| 2.6  | Evolution des différents critères de sélection du nombre de composantes en fonction du nombre de composantes pour l'exemple 1 . . . . .                                                                                                            | 26 |
| 2.7  | Evolution des différents critères de sélection du nombre de composantes en fonction du nombre de composantes pour l'exemple 2 . . . . .                                                                                                            | 27 |
| 2.8  | Différentes étapes pour la détermination du nombre de composantes (modèle ACP) et des variables à surveiller. . . . .                                                                                                                              | 28 |
| 2.9  | Evolution des mesures et des estimations des trois variables $x_1$ , $x_4$ et $x_7$ de l'exemple 1 . . . . .                                                                                                                                       | 30 |
| 2.10 | Evolution des mesures et des estimations des trois variables $x_1$ , $x_5$ et $x_9$ de l'exemple 2 . . . . .                                                                                                                                       | 31 |
| 3.1  | Décomposition du vecteur de mesure $\mathbf{x}$ en un vecteur de mesure estimé $\hat{\mathbf{x}}$ et un vecteur de résidus $\mathbf{e}$ . . . . .                                                                                                  | 37 |
| 3.2  | Interprétation géométrique de l'analyse en composantes principales . . . . .                                                                                                                                                                       | 38 |
| 3.3  | Illustration graphique de la condition de détectabilité avec le $SPE$ . . . . .                                                                                                                                                                    | 42 |
| 3.4  | Illustration graphique de la détection de défaut par les deux statistiques $SPE$ et $T^2$                                                                                                                                                          | 44 |
| 3.5  | Evolution du $SPE$ et du $SPE$ filtré avec un défaut affectant $x_3$ à partir de l'instant 300 (exemple 1) . . . . .                                                                                                                               | 48 |
| 3.6  | Evolution de la statistique $T^2$ avec un défaut affectant $x_3$ à partir de l'instant 300.                                                                                                                                                        | 48 |
| 3.7  | Evolution des $SPE$ filtrés correspondant aux différents défauts simulés et représentés sur le tableau (tab 3.1). . . . .                                                                                                                          | 49 |
| 3.8  | Evolution du $SPE$ et $SPE$ filtré en absence de défaut (cas de l'exemple 2) . . .                                                                                                                                                                 | 50 |
| 3.9  | Evolution du $SPE$ avec un défaut affectant la variable $x_3$ à partir de l'instant 300                                                                                                                                                            | 50 |
| 3.10 | Evolution du $\overline{SPE}$ avec un défaut affectant la variable $x_3$ à partir de l'instant 300                                                                                                                                                 | 51 |
| 3.11 | Evolution de l'indice $D_3$ filtré avec un défaut affectant la variable $x_3$ à partir de l'instant 300 . . . . .                                                                                                                                  | 53 |

|      |                                                                                                                                                                         |    |
|------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------|----|
| 3.12 | Evolution du $\overline{SPE}$ avec un défaut affectant la variable $x_9$ à partir de l'instant 350                                                                      | 53 |
| 3.13 | Evolution de l'indice $\bar{D}_1$ avec un défaut affectant la variable $x_9$ à partir de l'instant 350 . . . . .                                                        | 54 |
| 3.14 | Principe de la génération de résidus par estimation paramétrique dans le cas de l'ACP . . . . .                                                                         | 57 |
| 3.15 | Evolution des deux variables $y$ et $u$ en absence de défaut . . . . .                                                                                                  | 58 |
| 3.16 | Evolution de la variable $y$ avec un changement de 13% sur le paramètre $a$ à l'instant 1000. . . . .                                                                   | 58 |
| 3.17 | Evolution de $T^2$ et de $SPE$ en présence du défaut . . . . .                                                                                                          | 59 |
| 3.18 | Evolution de l'indice de détection $In_1$ en absence et en présence de défaut . . . .                                                                                   | 59 |
| 3.19 | Evolution des différents résidus filtrés des défauts affectant les différentes variables (exemple 1) . . . . .                                                          | 64 |
| 3.20 | Maximisation du produit $w_1^b b_2$ et $w_1^c$ est orthogonal à $b_1$ . . . . .                                                                                         | 66 |
| 3.21 | Evolutions des différents résidus structurés filtrés pour des défauts affectant les différentes variables (exemple 1) . . . . .                                         | 67 |
| 3.22 | Evolution des différents résidus structurés (filtrés avec un filtre EWMA, $\gamma = 0.1$ ) pour des défauts affectant les différentes variables (exemple 2) . . . . .   | 70 |
| 3.23 | Procédure de structuration de résidus par ACP partielles . . . . .                                                                                                      | 71 |
| 3.24 | Procédure de localisation par l'ACP partielle structurée . . . . .                                                                                                      | 71 |
| 3.25 | Evolutions des différents $SPE$ correspondant aux ACP partielles des différents défauts (exemple 1) . . . . .                                                           | 73 |
| 3.26 | Evolution des différents $SPE$ obtenus par la méthode d'élimination en absence de défauts. . . . .                                                                      | 74 |
| 3.27 | Evolution des différents $SPE$ obtenus par la méthode d'élimination avec un défaut affectant la variable $x_1$ à partir de l'instant 350 avec seuils modifiés . . . . . | 75 |
| 3.28 | Localisation de la variable $x_1$ avec le rapport $Q_r$ . . . . .                                                                                                       | 75 |
| 3.29 | Evolution des différents $D_{i,-j}$ et localisation de la variable $x_3$ . . . . .                                                                                      | 76 |
| 3.30 | Différents résidus obtenus à partir du modèle ACP de l'exemple 1. . . . .                                                                                               | 80 |
| 3.31 | Indices de validité de capteur filtrés avec un défaut affectant $x_1$ entre les instants 400 et 500 . . . . .                                                           | 81 |
| 3.32 | Evolution des différents indices $D_3$ filtrés après reconstruction des différents variables (exemple 2) . . . . .                                                      | 85 |
| 3.33 | Indices $A_3^{(j)}$ et localisation de la variable en défaut $x_3$ . . . . .                                                                                            | 86 |
| 3.34 | Evolution des différents indices $D_1$ filtrés après reconstruction des différents variables (exemple 2) . . . . .                                                      | 86 |
| 3.35 | Indices $A_1^{(j)}$ et localisation de la variable en défaut $x_9$ . . . . .                                                                                            | 87 |
| 3.36 | Contributions des variables à l'indice $\bar{D}_3$ avec un défaut affectant $x_3$ . . . . .                                                                             | 88 |
| 3.37 | Contributions des variables à l'indice $\bar{D}_1$ avec un défaut affectant $x_9$ . . . . .                                                                             | 88 |
| 3.38 | Evolution des deux variables $x_1$ et $x_2$ avec un défaut affectant $x_1$ à l'instant 430 et détecté à l'instant 434. . . . .                                          | 91 |
| 3.39 | Evolution de l'indice $D_4$ filtré avec un défaut affectant $x_1$ à l'instant 430. . . . .                                                                              | 91 |
| 3.40 | Contributions des différentes variables à l'indice $D_4$ à l'instant 434, approche de Kourti . . . . .                                                                  | 92 |
| 3.41 | Contributions des différentes variables à l'indice $D_4$ à l'instant 434, approche de Wise . . . . .                                                                    | 92 |
| 3.42 | Contributions des différentes variables à l'indice $D_4$ à l'instant 434, approche des variations . . . . .                                                             | 93 |
| 3.43 | Réorganisation de la matrice $\tilde{C}$ pour les défauts multiples . . . . .                                                                                           | 93 |

|      |                                                                                                                                                   |     |
|------|---------------------------------------------------------------------------------------------------------------------------------------------------|-----|
| 3.44 | Résultat de la localisation par reconstruction d'un défaut affectant les variables $x_3$ et $x_4$ .                                               | 95  |
| 3.45 | Résultat de la localisation par reconstruction de la variable $x_9$ avec un défaut affectant les variables $x_3$ et $x_9$ .                       | 96  |
| 3.46 | Résultat de la localisation par reconstruction de la variable $x_3$ avec un défaut affectant les variables $x_3$ et $x_9$ .                       | 97  |
| 4.1  | ACP linéaire.                                                                                                                                     | 103 |
| 4.2  | ACP non linéaire.                                                                                                                                 | 103 |
| 4.3  | Principe de la modélisation par l'analyse en composantes principales                                                                              | 104 |
| 4.4  | Projection des points sur la courbe.                                                                                                              | 106 |
| 4.5  | Utilisation d'un seul segment de droite                                                                                                           | 109 |
| 4.6  | Utilisation de deux segments de droite                                                                                                            | 109 |
| 4.7  | Utilisation de trois segments de droite                                                                                                           | 109 |
| 4.8  | Utilisation de quatres segments de droite                                                                                                         | 109 |
| 4.9  | Courbe principale pour l'exemple présenté avec l'algorithme de Verbeek en utilisant cinq segments de droite                                       | 110 |
| 4.10 | Réseau à cinq couches pour l'extraction d'une seule composante principale non linéaire                                                            | 111 |
| 4.11 | Mesures et estimation avec la première composante non linéaire avec un réseau à cinq couches                                                      | 114 |
| 4.12 | Extraction séquentielle des composantes principales non linéaires                                                                                 | 115 |
| 4.13 | Extraction simultanée des $\ell$ composantes principales non linéaires                                                                            | 116 |
| 4.14 | Réseau avec apprentissage des entrées (IT)                                                                                                        | 117 |
| 4.15 | Mesures et estimation de la courbe de l'exemple 1 en combinant les courbes principales et deux réseaux MLP a trois couches (jeu d'identification) | 118 |
| 4.16 | Mesures et estimation de la courbe de l'exemple 1 en combinant les courbes principales et deux réseaux MLP a trois couches (jeu de validation)    | 118 |
| 4.17 | Evolution de la première composante principale non-linéaire obtenue avec le réseau de neurones à cinq couches                                     | 118 |
| 4.18 | Evolution de la première composante principale non-linéaire obtenue avec l'algorithme de Verbeek                                                  | 118 |
| 4.19 | Réseau RBF pour la compression des données (projection)                                                                                           | 120 |
| 4.20 | Réseau RBF pour la décompression des données (projection inverse)                                                                                 | 120 |
| 4.21 | Réseau RBF pour le calcul de $\ell$ composantes principales                                                                                       | 121 |
| 4.22 | Mesures et estimation de la courbe en combinant les courbes principales et deux réseaux RBF à trois couches (jeu d'identification)                | 124 |
| 4.23 | Mesures et estimation de la courbe en combinant les courbes principales et deux réseaux RBF à trois couches (jeu de validation)                   | 124 |
| 4.24 | Mesure et estimation en utilisant une seule composante                                                                                            | 128 |
| 4.25 | Mesures et reconstruction de la courbe de l'exemple 2 avec un défaut sur $x_1$                                                                    | 129 |
| 4.26 | Erreur de reconstruction de la variable $x_1$                                                                                                     | 129 |
| 5.1  | Illustration du cycle de Chapman.                                                                                                                 | 133 |
| 5.2  | Implantation des capteurs, réseau AIRLOR                                                                                                          | 134 |
| 5.3  | Schéma du système d'acquisition et d'archivage des données caractérisant la qualité de l'air.                                                     | 135 |

|      |                                                                                                            |     |
|------|------------------------------------------------------------------------------------------------------------|-----|
| 5.4  | Concentration d'ozone des sites de Brabois, Dan et Fléville sur une période de quelques jours . . . . .    | 136 |
| 5.5  | Concentration de $NO_2$ des sites de Brabois, Dan et Fléville sur une période de quelques jours . . . . .  | 137 |
| 5.6  | Concentration de NO des sites de Brabois, Dan et Fléville sur une période de quelques jours . . . . .      | 137 |
| 5.7  | Concentration d'ozone et de $NO_2$ des sites de Fléville, Tomblaine et Lunéville. . .                      | 138 |
| 5.8  | Concentration d'ozone et de NO des sites de Fléville, Tomblaine et Lunéville. . .                          | 138 |
| 5.9  | Degrés de corrélation entre les capteurs (réseau AIRLOR). . . . .                                          | 139 |
| 5.10 | Evolution de la variance non reconstruite en fonction de $\ell$ . . . . .                                  | 142 |
| 5.11 | Mesure et estimation de l'ozone de Brabois ( $v_1$ ) par le modèle ACP linéaire. . . .                     | 143 |
| 5.12 | Mesure et estimation du $NO_2$ de Brabois ( $v_2$ ). . . . .                                               | 143 |
| 5.13 | Mesure et estimation du NO de Brabois ( $v_3$ ). . . . .                                                   | 144 |
| 5.14 | Evolution de l'erreur quadratique en présence d'un défaut affectant la variable 7. .                       | 145 |
| 5.15 | Evolution de l'indice $D_2$ en présence d'un défaut affectant la variable $v_7$ . . . . .                  | 145 |
| 5.16 | Localisation de la variable $v_7$ par reconstruction des 18 variables. . . . .                             | 146 |
| 5.17 | Localisation par l'indice $A_2^{(j)}$ calculé après la reconstruction de la $j^{eme}$ variable. .          | 147 |
| 5.18 | Localisation par calcul des contributions des variables à l'indice $D_2$ . . . . .                         | 147 |
| 5.19 | Reconstruction de la variable en défaut cas de l'ozone de Fléville $O_{3-FLE}$ . . . . .                   | 148 |
| 5.20 | Evolution de l'indice $D_3$ filtré en présence d'un défaut affectant la variable $v_{10}$ . .              | 149 |
| 5.21 | Localisation de la variable $v_{10}$ par reconstruction des 18 variables. . . . .                          | 150 |
| 5.22 | Localisation par l'indice $A_2^{(j)}$ calculé après la reconstruction de la $j^{eme}$ variable. .          | 151 |
| 5.23 | Localisation de la variable $v_{10}$ par calcul des contributions des variables à l'indice $D_3$ . . . . . | 151 |
| 5.24 | Reconstruction de la variable en défaut cas de l'ozone de St-Nicolas $O_{3-STN}$ . . .                     | 152 |
| 5.25 | Evolution de l'indice $D_1$ filtré en présence d'un défaut affectant la variable $v_2$ . . .               | 152 |
| 5.26 | Localisation de la variable $v_2$ par reconstruction des 18 variables. . . . .                             | 153 |
| 5.27 | Localisation par l'indice $A_2^{(j)}$ calculé après la reconstruction de la $j^{eme}$ variable. .          | 154 |
| 5.28 | Localisation de la variable $v_2$ par calcul des contributions des variables à l'indice $D_1$ . . . . .    | 154 |
| 5.29 | Reconstruction de la variable en défaut du dioxyde d'azote de Brabois $NO_{2-BRA}$ . .                     | 155 |
| 5.30 | Evolution de l'indice $D_1$ filtré en présence d'un défaut affectant la variable $v_3$ . . .               | 155 |
| 5.31 | Localisation de la variable $v_3$ par reconstruction des 18 variables. . . . .                             | 156 |
| 5.32 | Reconstruction de la variable en défaut $NO_{BRA}$ de Brabois. . . . .                                     | 157 |

# Notations

|                                                |                                                                                                           |
|------------------------------------------------|-----------------------------------------------------------------------------------------------------------|
| $X \in \mathbb{R}^{N \times m}$                | Matrice de données représentant le fonctionnement normal du système,                                      |
| $\hat{X}$                                      | Estimation de $X$ par le modèle ACP,                                                                      |
| $E$                                            | Matrice des résidus d'estimation de $X$ ,                                                                 |
| $\Sigma \in \mathbb{R}^{m \times m}$           | Matrice de covariance de $X$ ,                                                                            |
| $N$                                            | Nombre d'échantillons mesurés,                                                                            |
| $m$                                            | Nombre de variables (dimension de l'espace des données mesurées),                                         |
| $\ell$                                         | Nombre de composantes retenues dans le modèle ACP (dimension du sous-espace des composantes principales), |
| $k$                                            | Indice du temps,                                                                                          |
| $\mathbf{x} \in \mathbb{R}^m$                  | Nouveau vecteur de mesure,                                                                                |
| $\hat{\mathbf{x}}$                             | Estimation du vecteur $\mathbf{x}$ par le modèle ACP,                                                     |
| $\mathbf{x}_i \in \mathbb{R}^m$                | Nouveau vecteur de mesure dont la $i$ ème composante est reconstruite,                                    |
| $x_i$                                          | La $i$ ème composante du vecteur $\mathbf{x}$ ,                                                           |
| $\bar{\mathbf{x}}$                             | Vecteur moyen de $\mathbf{x}$ ,                                                                           |
| $\mathbf{x}^{(i)} \in \mathbb{R}^{m-1}$        | Le vecteur $\mathbf{x}$ sans la $i$ ème composante,                                                       |
| $P \in \mathbb{R}^{m \times m}$                | Matrice des vecteurs propres de $\Sigma$ ,                                                                |
| $\hat{P} \in \mathbb{R}^{m \times \ell}$       | Matrice des $\ell$ premiers vecteurs propres de $\Sigma$ ,                                                |
| $\tilde{P} \in \mathbb{R}^{m \times (m-\ell)}$ | Matrice des $m - \ell$ derniers vecteurs propres de $\Sigma$ ,                                            |
| $\hat{\mathbf{t}}$                             | Vecteurs des $\ell$ premières composantes principales,                                                    |
| $\tilde{\mathbf{t}}$                           | Vecteur des $m - \ell$ dernières composantes principales,                                                 |
| $\hat{C} = \hat{P}\hat{P}^T$                   | Matrice représentant le modèle ACP,                                                                       |
| $\lambda_i$                                    | $i$ ème valeur propre de $\Sigma$ ,                                                                       |
| $p_i$                                          | $i$ ème vecteur propre de $\Sigma$ correspondant à $\lambda_i$ ,                                          |
| $\mathcal{S}_{\hat{P}}$                        | Sous-espace des composantes principales,                                                                  |
| $\mathcal{S}_r$                                | Sous-espace des résidus,                                                                                  |
| $\mathcal{E}$                                  | Espérance mathématique,                                                                                   |
| $z_i$                                          | Valeur reconstruite de la mesure $x_i$ ,                                                                  |
| $\rho_i$                                       | Variance de l'erreur de reconstruction de la $i$ ème variable,                                            |
| $\xi_i$                                        | $i$ ème ligne d'une matrice identité $I_m$ ,                                                              |
| $e_i$                                          | Erreur d'estimation sur la $i$ ème variable,                                                              |
| $\mathbf{e}$                                   | Vecteur des erreurs d'estimation,                                                                         |
| $\bar{\mathbf{e}}$                             | Vecteur d'erreur filtré,                                                                                  |
| $\beta$                                        | Facteur d'oubli pour le filtrage des résidus,                                                             |
| $SPE$                                          | Erreur quadratique d'estimation (squared prediction error),                                               |
| $T^2$                                          | Statistique de Hotelling,                                                                                 |
| $\mathbf{r}$                                   | Vecteur de résidus structurés,                                                                            |
| $r_i$                                          | $i$ ème composante de $\mathbf{r}$ ,                                                                      |



# 1

## Introduction générale

Les enjeux économiques en constante évolution amènent à produire toujours plus. La moindre défaillance sur un processus est néfaste dans un environnement où le rendement est primordial. Il est donc nécessaire de s'assurer en permanence de la conduite optimale du procédé. L'information permettant de traduire le comportement d'un système est donnée par les mesures des variables de ce processus. La qualité des mesures est un élément essentiel pour permettre la surveillance et l'évaluation des performances d'un processus. La qualité de l'information peut être accrue en améliorant la précision de l'instrumentation et en multipliant les capteurs. Pour des raisons techniques ou financières, cette solution, où une même grandeur est mesurée par plusieurs capteurs, est réservée aux industries de haute technologie. De plus, cette redondance matérielle ne permet pas de se protéger contre une défaillance de certains éléments communs de la chaîne de mesure : plusieurs capteurs mesurant la même grandeur sont généralement géographiquement voisins et alimentés par le même réseau électrique ; une panne d'alimentation entraîne un arrêt de tout le système de mesure. L'exploitation de modèles a priori exacts liant différentes grandeurs mesurées offre un autre moyen pour vérifier la fiabilité des mesures. Cette redondance analytique présente l'avantage de ne pas augmenter le coût de l'installation et de se dégager des contraintes matérielles. Dans le domaine du diagnostic, des méthodes basées sur le concept de redondance de l'information ont été développées. Leur principe repose généralement sur un test de cohérence entre un comportement observé du processus fourni par des capteurs et un comportement prévu fourni par une représentation mathématique du processus. Les méthodes de redondance analytique nécessitent donc un modèle du système à surveiller. Ce modèle comprend un certain nombre de paramètres dont les valeurs sont supposées connues lors du fonctionnement normal. La comparaison entre le comportement réel du système et le comportement attendu donné par le modèle fournit une quantité, appelée résidu, qui va servir à déterminer si le système est dans un état défaillant ou non.

L'architecture générale du système de diagnostic à base de redondance analytique pour la détection et la localisation de défaillances de capteurs, sachant qu'on ne s'intéresse qu'à des défaillances d'instrumentation (capteurs ou actionneurs), est représentée sur la figure (fig 1.1).

Dans une première étape, il s'agit de comparer les observations avec les connaissances sur le comportement normal du système contenues dans un modèle afin de vérifier leur cohérence. Cette comparaison conduit à la génération d'indicateurs de défauts appelés résidus. Ces indicateurs

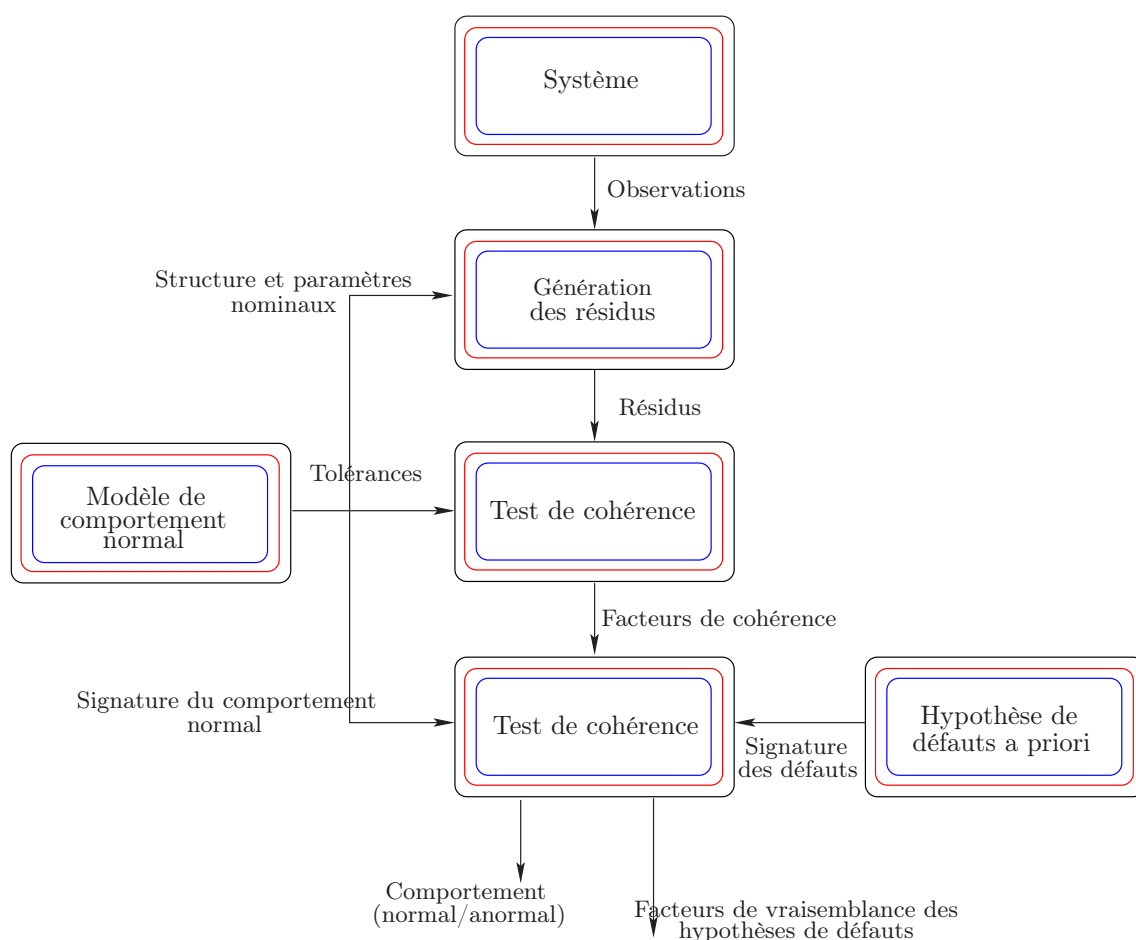


Figure 1.1 – Schéma de principe de la détection et de la localisation de défauts à base de modèles analytiques

sont souvent des écarts entre les caractéristiques observées et les caractéristiques de références qui définissent le comportement normal du système. Sur la (fig 1.3), on présente le principe de base de génération de résidus. Généralement, les méthodes de génération des résidus sont basées soit sur une estimation d'état ou une estimation paramétrique. Dans le cadre des méthodes reposant sur l'estimation d'état on retrouve trois approches :

1. l'approche de l'espace de parité : une représentation générale des différents aspects de cette méthode est donnée dans [6].
2. l'approche par observateur ou filtre de Kalman [20].
3. les filtres sensibles aux défauts (limités aux procédés invariants) [110].

En estimation paramétrique [45], le ou les vecteurs des paramètres sont estimés, à partir des mesures d'entrée et de sortie du système. Les paramètres du modèle (paramètres de référence) et les paramètres estimés sont alors comparés pour générer des résidus.

Les valeurs des résidus doivent refléter l'effet des défaillances. Elles doivent être proches de zéro en l'absence de défaut et différentes de zéro dans le cas contraire.

La seconde étape, appelée évaluation des résidus (fig 1.4), consiste à décider de la présence ou non d'anomalies de comportement au sein du système et à localiser la ou les composantes

dont le dysfonctionnement est à l'origine de ces anomalies. L'évaluation des résidus consiste à effectuer un premier test de cohérence, sous la forme d'une comparaison des résidus à des tolérances caractérisant le comportement normal du système, en prenant en compte leurs aspects non déterministes (bruits de mesure, erreurs de modélisation). Cette phase de détection permet alors de recenser les équations de modèle qui ne sont plus vérifiées révélant ainsi une ou plusieurs anomalies de comportement au sein du système. Les résidus sont donc conçus en vue de faciliter leur exploitation ultérieure par un outil de décision destiné à détecter et à localiser les défauts. Pour cela, deux approches sont possibles [23] :

- génération de résidus directionnels [24] : les résidus sont conçus de telle sorte que le vecteur des résidus reste confiné dans une direction particulière de l'espace des résidus, en réponse à un défaut particulier (fig 1.2).
- génération de résidus structurés [23] : les résidus sont conçus de façon à répondre à des sous-ensembles de défauts différents. Ces sous-ensembles de défauts permettent de structurer une table de signature appelée également matrice d'incidence ou matrice de signatures théoriques de défauts. Ces signatures traduisent l'influence des défauts sur les résidus. Pour que tous les défauts puissent être détectés, aucune colonne de la matrice des signatures théoriques de défauts ne doit être nulle, et pour que tous les défauts puissent être localisés, toutes les signatures théoriques doivent être distinctes. On peut distinguer trois cas pour une matrice d'incidence (tab 1.1) :
  - non localisante : une colonne est nulle ou deux au moins sont identiques.
  - faiblement localisante : les colonnes sont non nulles et distinctes deux à deux.
  - fortement localisante : en plus d'être faiblement localisante, aucune colonne ne peut être obtenue à partir d'une autre en remplaçant un 1 par un 0.

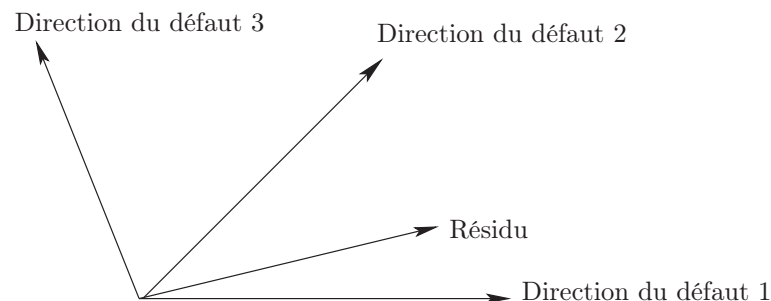


Figure 1.2 – Exemple de résidu directionnel ou le résidu est dans la direction du défaut 1

A l'issue de la phase de prise de décision, dans le cas des résidus structurés, un facteur de cohérence est obtenu pour chacune des équations de modèle. A chaque instant, l'ensemble des valeurs des facteurs de cohérence constitue ce que l'on appelle une signature expérimentale [22]. La phase de localisation consiste alors à effectuer un nouveau test de cohérence entre la signature expérimentale et les signatures de références afin de déterminer les hypothèses de défauts les plus vraisemblables.

Les méthodes de redondance analytique sont fondées sur l'utilisation des redondances présentes dans les mesures des variables. Le prix à payer est toutefois l'élaboration d'un modèle mathématique aussi complet que possible dont la qualité est primordiale pour l'obtention d'un système de détection performant. Cela est tout à fait envisageable lorsqu'il s'agit des systèmes de petite dimension. En revanche, pour les systèmes de taille plus importante, établir de telles relations mathématiques entre les variables paraît moins immédiat. De plus, le choix de modèle pour la génération de résidus est arbitraire et il se peut très bien que les corrélations entre certaines

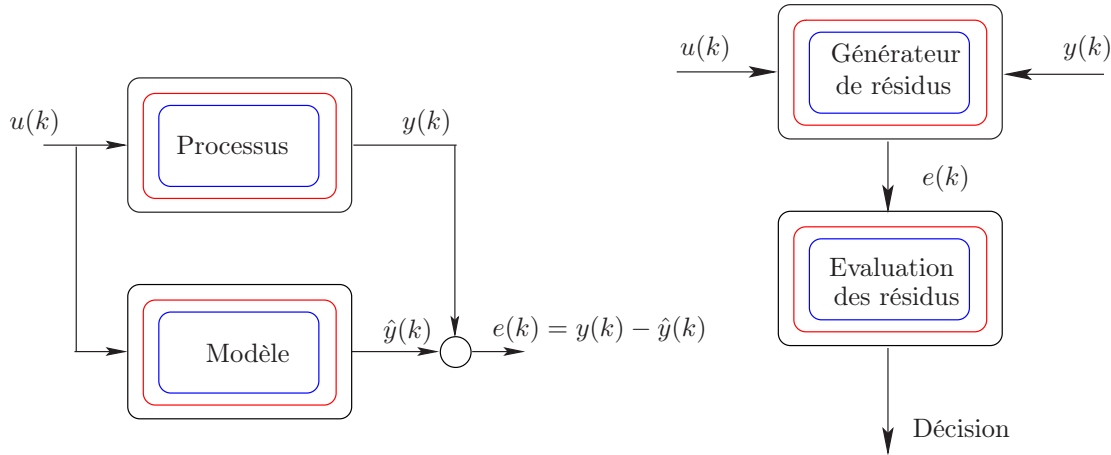


Figure 1.3 – Principe de génération des résidus

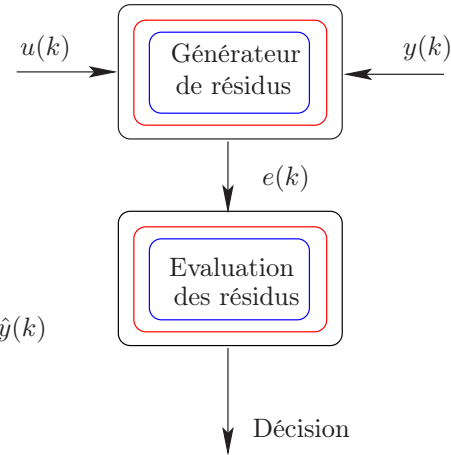


Figure 1.4 – Structure générale d'un système de détection

| ①     | $d_1$ | $d_2$ | $d_3$ | ②     | $d_1$ | $d_2$ | $d_3$ | ③     | $d_1$ | $d_2$ | $d_3$ |
|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|
| $r_1$ | 1     | 1     | 0     | $r_1$ | 1     | 1     | 0     | $r_1$ | 1     | 1     | 0     |
| $r_2$ | 1     | 1     | 1     | $r_2$ | 1     | 0     | 1     | $r_2$ | 1     | 0     | 1     |
| $r_3$ | 1     | 1     | 1     | $r_3$ | 1     | 1     | 1     | $r_3$ | 0     | 1     | 1     |

Table 1.1 – Exemple de structure de résidus. Les structures ② et ③ sont localisantes alors que ① est non localisante.

variables ne soient pas prises en compte. Comme alternative, les méthodes basées sur l'analyse en composantes principales (ACP), sont très intéressantes pour la mise en évidence des corrélations linéaires significatives entre les variables du processus sans formuler de façon explicite le modèle du système. Ainsi, toutes les corrélations entre les différentes variables sont prises en compte dans le modèle ACP. L'ACP est employée pour modéliser le comportement du processus en fonctionnement normal et les défauts sont alors détectés en comparant le comportement observé et celui donné par le modèle [57].

L'ACP permet de générer un modèle du processus basé sur la connaissance issue du système sans avoir une forme explicite d'un modèle entrées/sorties. Ainsi, elle permet d'exploiter toutes les relations linéaires qui peuvent exister entre les différentes variables. Dans ce mémoire nous allons étudier ce modèle particulier pour le diagnostic.

## Chapitre 2

Le deuxième chapitre présentera le principe de l'analyse en composantes principales (ACP). Les deux phases de modélisation et de génération de résidus avec l'ACP y seront exposées. L'identification du modèle ACP repose, généralement, sur deux étapes; la première consiste à déterminer la structure du modèle alors que la seconde consiste en l'estimation des paramètres. La deuxième étape est très simple et se ramène à un calcul de valeurs et vecteurs propres alors que la première, qui est du moins la plus difficile, consiste à déterminer le nombre de composantes principales à retenir dans le modèle ACP. Quelques critères classiques de sélection de ce nombre seront exposés dans ce chapitre. Un critère basé sur le principe de l'erreur de reconstruction, exploitant la redondance qui existe entre les variables, sera également présenté et adopté pour la

détermination de la structure du modèle ACP.

Une fois le modèle ACP identifié et les résidus générés, la phase de prise de décision pour la détection et la localisation de défauts fera l'objet du troisième chapitre.

### *Chapitre 3*

Ce chapitre sera consacré à la détection et localisation de défauts par analyse en composantes principales linéaires.

Pour que le lecteur puisse faire le lien avec des méthodes classiques, l'équivalence entre la génération de résidus avec les relations de parités et l'analyse en composantes principales est présentée.

Dans la plupart des travaux publiés utilisant l'analyse en composantes principales pour la détection de défauts, deux indices sont souvent utilisés. La statistique  $T^2$  de Hotelling calculée à partir des premières composantes principales et l'erreur quadratique d'estimation. Cependant, la statistique  $T^2$  ne représente pas vraiment un résidu puisqu'elle est calculée dans le sous espace principal (sous-espace représentant les variations significatives des variables du processus) et l'erreur quadratique d'estimation calculée dans le sous-espace résiduel (somme quadratique des résidus) est sensible aux erreurs de modélisation. L'objectif est de trouver un indice de détection permettant de s'affranchir de ce problème et d'être plus sensible aux défauts. Nous proposons alors un indice  $D_i$  utilisant les dernières composantes principales dans différents sous-espaces de l'espace résiduel [31, 32].

Pour la localisation de défaut, plusieurs approches seront exposées. La première approche est basée sur la structuration des résidus, comme dans le cas de l'espace de parité. D'autres approches utilisent le même principe que les approches classiques utilisant des bancs de modèles, comme l'approche par ACP partielles qui utilise des ACP avec des ensembles réduits de variables ou encore l'approche par élimination qui génère des résidus en éliminant une variable à chaque fois. Basée sur le même principe l'approche exploitant le principe de reconstruction sera également présentée. La dernière approche est une approche basée sur le calcul des contributions des variables à l'indice de détection [35].

Les deux dernières approches sont adaptées à l'indice de détection proposé pour la localisation de défauts.

Les deux exemples présentés dans le premier chapitre sont utilisés pour illustrer les différentes approches présentées dans ce chapitre.

### *Chapitre 4*

L'extension de l'analyse en composantes principales dans le cas non linéaire sera présentée dans ce chapitre.

Dans le cas non linéaire plusieurs approches seront présentées. L'approche des courbes principales a été présentée comme une technique de généralisation de l'ACP linéaire. Cependant cette approche ne permet pas d'obtenir un modèle de représentation permettant de calculer les composantes principales non linéaires en ligne à partir de nouvelles observations du processus à surveiller. Indépendamment de ces travaux, des représentations neuronales de l'analyse en composantes principales non linéaires (ACPNL) ont été développées. Nous présenterons les principales représentations et leurs principes.

Dans le cas de l'ACPNL utilisant les réseaux de fonctions à base radiale RBF (Radial Basis Function), nous proposons de combiner ce type de réseau avec les courbes principales. Ainsi, le problème se ramène à un problème de régression linéaire avec un gain considérable en temps de calcul.

Toutefois, il est important de déterminer le nombre de composantes à retenir dans le modèle. Pour cette raison, nous proposons une extension du critère de la variance non reconstruite dans le cas non linéaire [36].

Un exemple de simulation sera présenté tout au long de ce chapitre pour illustrer les différentes approches présentées.

## **Chapitre 5**

Le dernier chapitre de cette thèse sera consacré à l'application de l'analyse en composantes principales pour la détection et la localisation de défauts de capteurs d'un réseau de surveillance de la qualité de l'air en Lorraine [31, 32].

Beaucoup d'activités humaines produisent des polluants primaires comme les oxydes d'azote ( $\text{NO}_2$  et  $\text{NO}$ ), le dioxyde de soufre et les composés organiques volatiles (COV) qui forment dans la basse atmosphère, par des réactions chimiques ou photochimiques, des polluants secondaires comme l'ozone. Un certain nombre de ces polluants sont susceptibles de poser des problèmes pour la santé humaine et les systèmes écologiques.

L'étude présentée concerne trois polluants qui sont l'ozone et les oxydes d'azote  $\text{NO}_2$  et  $\text{NO}$  mesurés sur différents sites en Lorraine.

Un modèle ACP est donc calculé à partir des mesures disponibles et la détection de défauts est effectuée en utilisant l'indice de détection proposé dans le deuxième chapitre. La variable en défaut est ainsi localisée en utilisant l'approche par reconstruction ou l'approche des contributions appliquées à l'indice de détection proposé. Les résultats obtenus sont discutés.

## 2

## Analyse en Composantes Principales Linéaires

## Sommaire

|            |                                                          |           |
|------------|----------------------------------------------------------|-----------|
| <b>2.1</b> | <b>Introduction</b>                                      | <b>7</b>  |
| <b>2.2</b> | <b>Principes de l'analyse en composantes principales</b> | <b>9</b>  |
| <b>2.3</b> | <b>Identification du modèle ACP</b>                      | <b>14</b> |
| 2.3.1      | Estimation des paramètres du modèle                      | 14        |
| 2.3.2      | Détermination de la structure du modèle                  | 18        |
| 2.3.3      | Sélection de capteurs                                    | 25        |
| <b>2.4</b> | <b>Conclusion</b>                                        | <b>28</b> |

## 2.1 Introduction

Dans le domaine du diagnostic, des méthodes basées sur le concept de redondance de l'information ont été développées. Leur principe repose généralement sur un test de cohérence entre un comportement observé du processus fourni par des capteurs et un comportement prévu fourni par une représentation mathématique du processus.

Les méthodes de redondance analytique nécessitent un modèle du système à surveiller. Ce modèle comprend un certain nombre de paramètres dont les valeurs sont supposées connues lors du fonctionnement normal. La comparaison entre le comportement réel du système et le comportement attendu donné par le modèle fournit une quantité appelée résidu qui va servir à déterminer si le système est dans un état défaillant ou non.

Les méthodes statistiques utilisées pour la détection de défauts dans des processus industriels sont incluses dans un domaine généralement connu sous le nom de maîtrise ou contrôle statistique de processus (CSP) (traduction française de SPC, Statistical Process Control). Les origines du CSP remontent aux travaux de Shewhart publiés en 1931 [85]. Les techniques les plus utilisées et les plus populaire du CSP sont des méthodes univariées, c'est-à-dire, surveillant chaque variable ainsi que ses statistiques (moyenne et variance) indépendamment des autres variables.

Cette approche exige un opérateur surveillant sans interruption peut être des douzaines de différents diagrammes univariés, ce qui réduit sensiblement sa capacité à faire des évaluations

précises au sujet de l'état du processus. Ainsi, l'application de cette approche à des systèmes de grande dimension (grand nombre de variable) devient difficile sinon impossible.

De plus le calcul de la moyenne et de l'écart-type donne, pour chaque variable, des informations concernant l'ordre de grandeur et la dispersion des données ; le calcul de la matrice de corrélation des variables donne des indications sur l'évolution simultanée des variables prises deux à deux. Ces éléments de statistique descriptive univariée et bivariée ne donnent cependant aucune information lorsque les variables sont considérées simultanément. Cette étude simultanée des variables est précisément le but de l'analyse en composantes principales.

L'utilisation de l'analyse en composantes principales pour l'exploitation des données remonte au début du siècle dernier. Elle est principalement issue des travaux de psychomètres américains [75, 91, 41, 95].

L'analyse en composantes principales (ACP) est une technique descriptive permettant d'étudier les relations qui existent entre les variables, sans tenir compte, a priori, d'une quelconque structure [48, 11]. Le but de l'ACP est d'identifier la structure de dépendance entre des observations multivariées afin d'obtenir une description ou une représentation compacte de ces dernières. Son utilisation a été restreinte à la projection des données sur les différents axes factoriels et au calcul de distances par rapport à ces axes comme outil de détection de valeurs aberrantes. L'utilisation de l'ACP de cette manière n'est pas très pratique, car un opérateur doit visualiser les projections pour prendre une décision quant à la présence de valeurs aberrantes.

Depuis les années 70, de nombreux travaux ont proposé d'utiliser l'analyse en composantes principales comme un outil de modélisation des processus à partir de laquelle un modèle peut être obtenu [48, 114, 57, 62]. Ainsi, cette alternative permet d'estimer les variables ou les paramètres du processus à surveiller.

Mathématiquement l'ACP est une technique de projection orthogonale linéaire qui projette les observations multidimensionnelles représentées dans un sous-espace de dimension  $m$  ( $m$  est le nombre de variables observées) dans un sous-espace de dimension inférieure ( $\ell < m$ ) en maximisant la variance des projections.

La solution de ce problème de maximisation définit à la fois la projection du sous-espace de dimension  $m$  dans le sous-espace de dimension  $\ell$  et la projection inverse (du sous-espace de dimension  $\ell$  vers le sous-espace de dimension  $m$ ) permettant d'estimer les variables originelles.

Dans ce sens, l'ACP peut être considérée comme une technique de minimisation de l'erreur quadratique d'estimation ou une technique de maximisation de la variance des projections (il faut noter que ces deux critères sont équivalents).

Dans ce travail l'ACP est utilisée comme un outil de modélisation des relations linéaires entre les différentes grandeurs représentant le comportement d'un processus quelconque.

L'estimation des paramètres du modèle ACP est effectué par calcul des valeurs et vecteurs propres de la matrice de corrélation des données. Cependant, pour la détermination de la structure du modèle, il faut déterminer le nombre de composantes à retenir dans ce modèle. Pour cette raison, plusieurs critères de sélection du nombre de composantes seront présentés.

Le critère basé sur le principe de la variance de reconstruction est adopté pour la sélection du nombre de composantes à retenir dans le modèle ACP car il exploite les redondances entre les variables. Ce nombre permettra d'identifier le modèle ACP dont dépend étroitement la procédure de détection et de localisation.

Il faut noter que nous avons travaillé uniquement sur un modèle ACP statique et que l'extension au cas dynamique est possible [107, 78].

## 2.2 Principes de l'analyse en composantes principales

Considérons un vecteur de données aléatoires  $\mathbf{x}(k) = [x_1, \dots, x_m]^T \in \mathbb{R}^m$  de moyenne nulle  $\mathcal{E}\{\mathbf{x}(k)\} = 0$  et de matrice de covariance ou d'auto-corrélation  $\Sigma = \mathcal{E}\{\mathbf{x}\mathbf{x}^T\} \in \mathbb{R}^{m \times m}$ .

En analyse en composantes principales, un vecteur caractéristique  $\mathbf{t} \in \mathbb{R}^\ell$  est associé à chaque vecteur de données dont il optimise la représentation au sens de la minimisation de l'erreur d'estimation de  $\mathbf{x}$  ou la maximisation de la variance de  $\mathbf{t}$ . Les vecteurs  $\mathbf{t}$  et  $\mathbf{x}$  sont liés par une transformation linéaire  $\mathbf{t} = P^T \mathbf{x}$ , où la matrice de transformation  $P \in \mathbb{R}^{m \times \ell}$  vérifie la condition d'orthogonalité  $P^T P = I_\ell$ . Les colonnes de la matrice  $P$  forment les vecteurs de base orthonormés d'un sous-espace  $\mathbb{R}^\ell$  de représentation réduite des données. La transformation linéaire s'apparente ainsi à une projection de l'espace des données de dimension  $m$  vers un sous-espace orthogonal de dimension  $\ell$ .

Les composantes  $t_j$ , avec  $(j = 1, \dots, \ell)$ , du vecteur caractéristique  $\mathbf{t}$  représentent les composantes projetées du vecteur de données  $\mathbf{x}$  dans ce sous-espace.

Au sens de l'ACP, la projection  $P$  est optimale si l'erreur quadratique d'estimation des vecteurs de données  $\mathbf{x}$  est minimale. Ce problème d'optimisation s'exprime par :

$$P_{opt} = \arg \min_P J_e(P) \quad (2.1)$$

où  $J_e$  représente le critère d'erreur d'estimation de l'ACP. Sous la contrainte d'orthogonalité de la matrice de projection  $P^T P = I_\ell$ , ce critère aura la forme :

$$\begin{aligned} J_e(P) &= \mathcal{E}\{\|\mathbf{x} - \hat{\mathbf{x}}\|^2\} = \mathcal{E}\{\|\mathbf{x} - PP^T \mathbf{x}\|^2\} \\ &= \mathcal{E}\{(\mathbf{x} - P\mathbf{t})^T (\mathbf{x} - P\mathbf{t})\} \\ &= \mathcal{E}\{\mathbf{x}^T \mathbf{x} - 2\mathbf{t}^T \mathbf{t} + \mathbf{t}^T P P^T \mathbf{t}\} \\ &= \mathcal{E}\{trace(\mathbf{x}\mathbf{x}^T) - \mathbf{t}^T \mathbf{t}\} = trace(\Sigma) - \mathcal{E}\{\mathbf{t}^T \mathbf{t}\} \end{aligned} \quad (2.2)$$

Notons que la trace d'une matrice carrée est définie par la somme de ces éléments diagonaux. Du fait que la matrice de covariance  $\Sigma$  est indépendante de la matrice des paramètres  $P$ , minimiser  $J_e$  revient à maximiser le second terme  $J_v$  de son expression :

$$J_v(P) = \mathcal{E}\{\mathbf{t}^T \mathbf{t}\} = \sum_{j=1}^{\ell} \mathcal{E}\{t_j^2\} \quad (2.3)$$

Ainsi, la minimisation de l'erreur quadratique d'estimation de  $\mathbf{x}$  est équivalente à la maximisation de la variance des projections  $t_j$  des données. En conséquence :

$$P_{opt} = \arg \min_P J_e(P) = \arg \max_P J_v(P) \quad (2.4)$$

Le problème de l'ACP, considéré sous l'angle de la maximisation de la variance de projection des données, est celui de la détermination des vecteurs propres de la matrice de covariance  $\Sigma$ .

Afin d'établir ce fait, considérons un vecteur de données  $\mathbf{x} \in \mathbb{R}^m$  centré ( $\mathcal{E}\{\mathbf{x}\} = 0$ ) et déterminons la direction du vecteur directeur  $p \in \mathbb{R}^m$  unitaire ( $\|p\| = 1$ ) suivant laquelle la variance de la projection de  $\mathbf{x}$  est maximale. La projection  $t \in \mathbb{R}$  du vecteur de données  $\mathbf{x}$  le long de la direction représentée par le vecteur unitaire  $p$  est décrite par le produit scalaire  $t = \mathbf{x}^T p = p^T \mathbf{x}$  sous la contrainte  $\|p\|^2 = p^T p = 1$ . Cette projection constitue une variable dont la moyenne et la variance dépendent des propriétés statistiques des données. Sous l'hypothèse de nullité de la moyenne du vecteur de données  $\mathbf{x}$ , la valeur moyenne de la projection  $t$  est également

nulle :  $\mathcal{E}\{t\} = p^T \mathcal{E}\{\mathbf{x}\} = 0$ . En conséquence, la variance de la projection,  $\text{var}\{t\}$ , s'identifie à sa valeur quadratique :

$$\begin{aligned} \text{var}\{t\} &= \mathcal{E}\{(t - \mathcal{E}\{t\})^2\} = \mathcal{E}\{t^2\} \\ &= \mathcal{E}\{(p^T \mathbf{x})(\mathbf{x}^T p)\} \\ &= p^T \mathcal{E}\{\mathbf{x}\mathbf{x}^T\} p = p^T \Sigma p \end{aligned} \quad (2.5)$$

**Remarque 2.2.1** Lorsque le vecteur de données présente une moyenne non nulle, celle-ci lui est préalablement retranchée avant analyse.

La maximisation de l'expression précédente de la variance de la projection, sous condition de norme unité du vecteur  $p$ , est un problème d'optimisation sous contrainte égalité qui peut être résolu au moyen de la méthode des multiplicateurs de Lagrange. Sous cette dernière forme, le problème d'optimisation sous contrainte égalité maximisant la variance de la projection est formalisé par la fonction de Lagrange :

$$\mathcal{L}(p, \lambda) = J_v(p) - \lambda (p^T p - 1) = p^T \Sigma p - \lambda (p^T p - 1) \quad (2.6)$$

où  $\lambda \in \mathbb{R}$  désigne le multiplicateur de Lagrange. Le vecteur  $p$  minimisant le critère d'optimisation (2.6) est alors solution du système d'équations suivant :

$$\begin{cases} \partial \mathcal{L}(p, \lambda) / \partial p = 0 \\ \partial \mathcal{L}(p, \lambda) / \partial \lambda = 0 \end{cases} \quad (2.7)$$

En utilisant la propriété de symétrie de la matrice de covariance des données  $\Sigma^T = \Sigma$ , le système d'équations s'écrit :

$$\begin{cases} \Sigma p - \lambda p = 0 \\ p^T p - 1 = 0 \end{cases} \quad (2.8)$$

La résolution de ce système d'équations est identifiée comme un problème d'estimation de valeur propre et de vecteur propre de la matrice de covariance sous contrainte de normalisation du vecteur propre. Ce problème est communément rencontré en algèbre linéaire. L'équation (2.8) admet des solutions réelles de la variable  $\lambda$  obtenues par résolution directe de l'équation caractéristique :

$$\text{Det}(\Sigma - \lambda I_m) = 0 \quad (2.9)$$

où  $\text{Det}(\cdot)$  est le déterminant d'une matrice carrée. Ces solutions sont appelées valeurs propres de la matrice de covariance  $\Sigma$ . A ces valeurs propres sont associés  $m$  vecteurs caractéristiques  $p$  appelés vecteurs propres. Tout vecteur propre  $p$  associé à une valeur propre  $\lambda$  est solution non triviale de l'équation :

$$(\Sigma - \lambda I_m) p = 0 \quad (2.10)$$

Notons par  $\lambda_1, \dots, \lambda_m$  les  $m$  valeurs propres de la matrice de covariance  $\Sigma$  et par  $p_1 \in \mathbb{R}^m, \dots, p_m \in \mathbb{R}^m$  les  $m$  vecteurs propres qui leurs sont associés. Nous pouvons alors écrire :

$$\Sigma p_i = p_i \lambda_i, \quad i = 1, \dots, m \quad (2.11)$$

ou sous forme matricielle

$$\Sigma P = P \Lambda \quad (2.12)$$

avec  $P = [p_1, p_2, \dots, p_m] \in \mathbb{R}^{m \times m}$  et  $\Lambda = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_m) \in \mathbb{R}^{m \times m}$ . La notation  $\text{diag}(\lambda_1, \lambda_2, \dots, \lambda_m)$  se réfère à la matrice carrée dont les seuls éléments non nuls situés sur la diagonale sont les valeurs  $\lambda_1, \dots, \lambda_m$ . En raison de la propriété de symétrie de la matrice de covariance  $\Sigma$ , les  $m$  valeurs propres  $\lambda_i$  sont réelles et les  $m$  vecteurs propres  $p_i$  sont distincts et orthogonaux. Si l'on ajoute à cette propriété la contrainte de norme unité, les vecteurs propres  $p_i$  ( $i = 1, \dots, m$ ) forment une base orthonormée :

$$P^T P = P P^T = I_m \quad (2.13)$$

Par définition, la transposée de la matrice carrée orthogonale correspond à son inverse, c'est-à-dire :  $P^T = P^{-1}$ . En conséquence, l'équation matricielle (2.12) admet la forme équivalente :

$$P^T \Sigma P = \Lambda \quad (2.14)$$

qui s'écrit sous forme développée :

$$p_i^T \Sigma p_j = \begin{cases} \lambda_i & \text{si } j = i \\ 0 & \text{si } j \neq i \end{cases} \quad (2.15)$$

De la comparaison des relations (2.5) et (2.15), il résulte :

$$\text{var}\{t_i\} = \mathcal{E}\{t_i^2\} = p_i^T \Sigma p_i = \lambda_i, i = 1, \dots, m \quad (2.16)$$

L'équation (2.16) révèle que les valeurs propres de la matrice de covariance  $\Sigma$  représentent les variances des projections  $t_i$  des données sur les directions représentées par les vecteurs propres  $p_i$  ( $i = 1, \dots, m$ ). En conclusion, la direction suivant laquelle la variance de la projection du vecteur de données  $\mathbf{x}$  est maximale, est représentée par le vecteur propre  $p_i$  correspondant à la valeur propre maximale  $\lambda_i$ . Le second axe factoriel rend la variance maximale tout en étant orthogonal au premier. De façon plus générale, le sous-espace vectoriel de dimension  $\ell$  qui assure une dispersion maximale des observations est défini par une base orthonormée formée des  $\ell$  vecteurs propres correspondant aux  $\ell$  plus grandes valeurs propres de la matrice  $\Sigma$ . Une décomposition en valeurs propres de la matrice  $\Sigma$ ,  $\Sigma = P \Lambda P^T$  révèle la structure de la matrice de covariance.

La projection  $\mathbf{t} = P^T \mathbf{x}$  d'un vecteur d'observation  $\mathbf{x} \in \mathbb{R}^m$  est le nouveau vecteur des variables transformées qui sont indépendantes et dont les variances sont les valeurs propres de la matrice  $\Lambda$  contenant les valeurs propres de  $\Sigma$  ordonnées dans l'ordre décroissant.

Les  $m$  vecteurs propres unitaires  $p_i$  de la matrice de covariance  $\Sigma$  représentent les  $m$  directions orthogonales de l'espace des données suivant lesquelles les variances des projections  $t_i$  des données sont maximales (fig 2.1). Ces directions englobent donc l'ensemble de l'information véhiculée par les données. Aussi, l'espace engendré par les vecteurs propres constitue-t-il un espace de représentation optimale des données. Dans cet espace, les composantes principales  $t_i$  du vecteur de données  $\mathbf{x}$  sont définies par :

$$t_i = p_i^T \mathbf{x} = \mathbf{x}^T p_i \quad i = 1, \dots, m \quad (2.17)$$

Celles-ci sont dénommées composantes principales et sont statistiquement non corrélées. En effet, en vertu des relations vectorielles

$$p_i^T \Sigma p_j = p_i^T \mathcal{E}\{\mathbf{x} \mathbf{x}^T\} p_j = \mathcal{E}\{p_i^T \mathbf{x} \mathbf{x}^T p_j\} = \mathcal{E}\{t_i t_j\} = 0 \quad i \neq j \quad (2.18)$$

La transposition matricielle des relations de projection de type (2.17) fournit l'expression analytique de l'analyse en composantes principales du vecteur de données  $\mathbf{x}$  :

$$\mathbf{t} = P^T \mathbf{x} \quad (2.19)$$

où  $P^T$  représente la matrice de projection optimale des données au sens de l'analyse en composantes principales.

Rappelons que tout vecteur de données  $\mathbf{x}$  peut être représenté par la combinaison linéaire des  $m$  vecteurs propres  $p_i$  ( $i = 1, \dots, m$ ) de la matrice de covariance  $\Sigma$ , pondérés par les composantes principales  $t_i = p_i^T \mathbf{x}$ . L'estimation d'un vecteur de données  $\mathbf{x}$  à partir de son vecteur de composantes principales associé  $\mathbf{t}$  est triviale. Il suffit pour cela de multiplier à droite chacun des membres de l'équation (2.19) par  $P$ , il en découle :

$$\mathbf{x} = P\mathbf{t} = \sum_{i=1}^m t_i p_i \quad (2.20)$$

L'intérêt de l'analyse en composantes principales réside dans la réduction de dimension de représentation de données ou plus simplement la compression de données. En effet, cette technique permet de caractériser les directions orthogonales d'un espace de données porteuses du maximum d'information au sens de la maximisation des variances de projections. L'amplitude des valeurs propres de la matrice de covariance  $\Sigma$  des données quantifie pour chacune de ces directions la quantité d'information encodée.

Rappelons que ces valeurs propres,  $\lambda_1, \dots, \lambda_m$ , sont non négatives par définition. La direction de l'espace des données matérialisée par le vecteur propre  $p_1$  associé à la plus grande valeur propre  $\lambda_1$  est la plus riche d'information. Inversement, la direction du vecteur propre directeur  $p_m$  associé à la plus petite valeur propre  $\lambda_m$  est celle qui présente le minimum d'information.

Il est donc possible de réduire la dimension de la représentation des données en ne retenant de l'expression (2.20) que les termes  $t_j p_j$  ( $j = 1, \dots, \ell$ ) associés aux  $\ell$  plus grandes valeurs propres  $\lambda_j$ . L'estimation  $\hat{\mathbf{x}}$  d'un vecteur de données  $\mathbf{x}$  est alors décrite par l'expression réduite :

$$\hat{\mathbf{x}} = \sum_{j=1}^{\ell} t_j p_j = \sum_{j=1}^{\ell} (p_j^T \mathbf{x}) p_j \quad (2.21)$$

Les données sont ainsi encodées par l'intermédiaire des  $\ell$  composantes principales  $t_1, \dots, t_\ell$  présentant les plus fortes variances, en comparaison des  $m$  valeurs descriptives  $x_1, \dots, x_m$  initialement requises.

La perte d'information induite par la réduction de dimension de représentation de chaque vecteur de données  $\mathbf{x}$  est mesurée par la différence  $\mathbf{e}$  entre ses représentations exacte (2.20) et approchée (2.21) :

$$\mathbf{e} = \mathbf{x} - \hat{\mathbf{x}} = \sum_{i=\ell+1}^m t_i p_i \quad (2.22)$$

Les  $(m - \ell)$  composantes principales  $t_i$  ( $i = \ell + 1, \dots, m$ ) à partir desquelles l'erreur d'estimation  $\mathbf{e}$  est évaluée, sont associées aux plus faibles valeurs propres  $\lambda_{\ell+1}, \dots, \lambda_m$ . Il est par conséquent bien évident que la compression de données préserve d'autant mieux l'information qu'elles véhiculent que ces valeurs propres sont de faibles valeurs. La somme des  $m - \ell$  valeurs propres minimales quantifie par ailleurs la perte d'information :

$$\mathcal{E}\{\mathbf{e}^T \mathbf{e}\} = \sum_{i=\ell+1}^m \text{var}\{t_i\} \quad (2.23)$$

Ainsi, le vecteur de données  $\mathbf{x}$  peut être exprimé sous la forme :

$$\mathbf{x} = \hat{\mathbf{x}} + \mathbf{e} \quad (2.24)$$

Dans la suite nous allons représenter la matrice des  $\ell$  premiers vecteurs propres par  $\hat{P}$ , d'où :

$$\hat{\mathbf{t}} = \hat{P}^T \mathbf{x} \quad (2.25)$$

où  $\hat{\mathbf{t}}$  est le vecteur des premières composantes principales et

$$\hat{\mathbf{x}} = \hat{C} \mathbf{x} \quad (2.26)$$

où  $\hat{C} = \hat{P} \hat{P}^T$ .

Ainsi, l'erreur quadratique d'estimation est donnée par :

$$SPE = \mathbf{e}^T \mathbf{e} \quad (2.27)$$

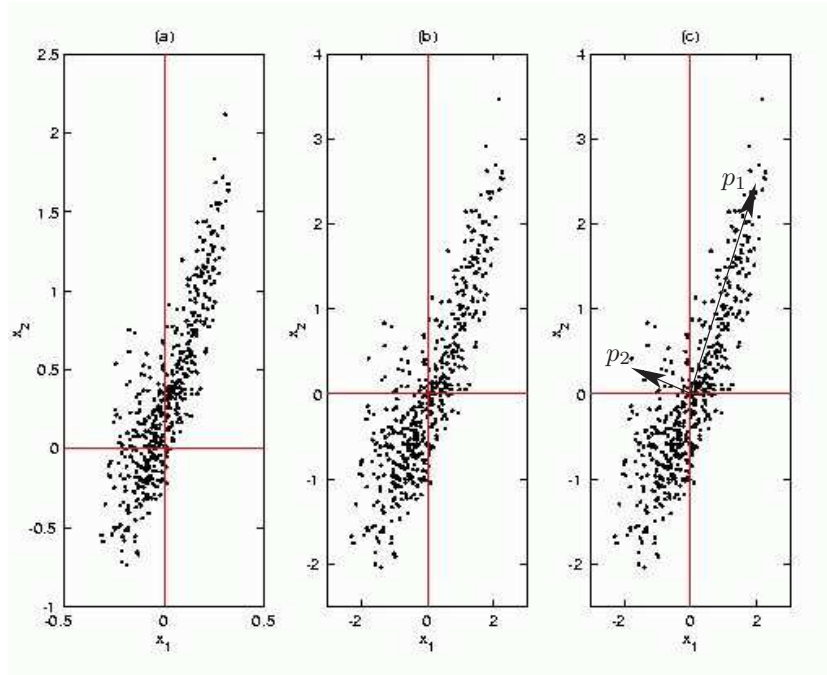


Figure 2.1 – Déroulement d'une analyse en composantes principales. (a) Distribution d'entrée. (b) Centrage et réduction de cette distribution. (c) Les deux axes principaux, correspondant aux vecteurs propres de la matrice de covariance de la distribution.

Ainsi, l'analyse en composantes principales peut être considérée comme une approche de modélisation avec laquelle on peut obtenir un modèle du système. On peut montrer également que les derniers vecteurs propres représentent en fait les relations linéaires ou quasi linéaires qui existent entre les variables [44]. Cette solution peut être obtenue par d'autres méthodes classique comme la méthode TLS (Total Least Square) [100].

## 2.3 Identification du modèle ACP

En pratique, pour la modélisation d'un processus en utilisant l'analyse en composantes principales, les mesures des variables du processus représentant tous les modes en fonctionnement normal de ce dernier sont collectées dans une matrice  $X^b$ . Soit  $m$  le nombre de variables et  $N$  le nombre d'observations de chaque variable.  $X^b \in \mathbb{R}^{N \times m}$  est donnée par :

$$X^b = \begin{pmatrix} x_1(1) & x_2(1) & \cdot & \cdot & \cdot & x_m(1) \\ x_1(2) & x_2(2) & \cdot & \cdot & \cdot & x_m(2) \\ \cdot & \cdot & & & & \cdot \\ \cdot & \cdot & & & & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ x_1(N) & x_2(N) & \cdot & \cdot & \cdot & x_m(N) \end{pmatrix} \quad (2.28)$$

où  $x_1(1)$  représente la mesure de la première variable du premier instant de mesure. Au préalable, afin de rendre le résultat indépendant des unités utilisées pour chaque variable, un pré-traitement indispensable consiste à centrer et réduire les variables. Chaque colonne  $X_j$  de la nouvelle matrice centrée est donnée par :

$$X_j = \frac{X_j^b - M_j}{\sigma_j} \quad (2.29)$$

où  $X_j^b$  est la  $j^{ieme}$  colonne de la matrice  $X^b$  et  $M_j$  est sa moyenne donnée par :

$$M_j = \frac{1}{N} \sum_{k=1}^N x_j(k) \quad (2.30)$$

et  $\sigma_j^2$  est sa variance qui sera estimée en utilisant l'équation :

$$\sigma_j^2 = \frac{1}{N} \sum_{k=1}^N (x_j(k) - M_j)^2 \quad (2.31)$$

La nouvelle matrice des données normalisées est notée :

$$X = [X_1 \quad \cdot \quad \cdot \quad \cdot \quad X_m] \quad (2.32)$$

La matrice de corrélation est donnée par :

$$\Sigma = \frac{1}{N-1} X^T X \quad (2.33)$$

Généralement, la procédure d'identification de modèles, consiste, après le choix d'une classe de modèle, à choisir une structure fixe puis à estimer les paramètres du modèle et enfin à valider ce modèle. Dans le cas de l'ACP, l'estimation des paramètres du modèle est très simple et revient en fait à un calcul de valeurs et vecteurs propres. Cependant le choix de la structure est plus délicat comme nous le verrons par la suite.

### 2.3.1 Estimation des paramètres du modèle

L'estimation des paramètres du modèle ACP se résume en une estimation des valeurs et vecteurs propres de la matrice de corrélation  $\Sigma$ . Une décomposition spectrale de cette dernière permet d'écrire :

$$\Sigma = P\Lambda P^T = \sum_{i=1}^m \lambda_i p_i p_i^T \quad (2.34)$$

où  $p_i$  est le  $i^{eme}$  vecteur propre de  $\Sigma$  et  $\lambda_i$  est la valeur propre correspondante.

S'il existe  $q$  relations linéaires entre les colonnes de  $X$ , on aura  $q$  valeurs propres nulles et la matrice  $X$  peut être représentée par les premières  $(m - q) = \ell$  composantes principales correspondant aux valeurs propres non nulles. Toutefois, les valeurs propres égales à zéro sont rarement rencontrées en pratique (relations quasi-linéaires, bruits,...). Donc, il est nécessaire de déterminer le nombre  $\ell$  représentant le nombre de vecteurs propres correspondant aux valeurs propres dominantes.

Pour illustrer ce qui a été dit jusqu'à présent sur l'ACP linéaire, nous présentons ici un premier exemple de simulation. Cet exemple sera également utilisé pour illustrer les différentes méthodes présentées dans ce premier chapitre ainsi que le deuxième chapitre.

Nous disposons de 7 variables représentant un système statique et qui sont décrites par les équations suivantes :

$$\begin{aligned} x_1 &= u_1 + \varepsilon_1 \\ x_2 &= u_1 + \varepsilon_2 \\ x_3 &= u_1 + \varepsilon_3 \\ x_4 &= u_2 + \varepsilon_4 \\ x_5 &= u_2 + \varepsilon_5 \\ x_6 &= 3u_1 + 2u_2 + \varepsilon_6 \\ x_7 &= 2u_1 + u_2 + \varepsilon_7 \end{aligned} \quad (2.35)$$

où les bruits de mesure  $\varepsilon_i$  sont des bruits aléatoires uniformément réparties entre  $-0.05$  et  $+0.05$ ,  $u_1$  et  $u_2$  sont des signaux en forme de crêteaux dont les amplitudes et les durées changent de manière aléatoire.

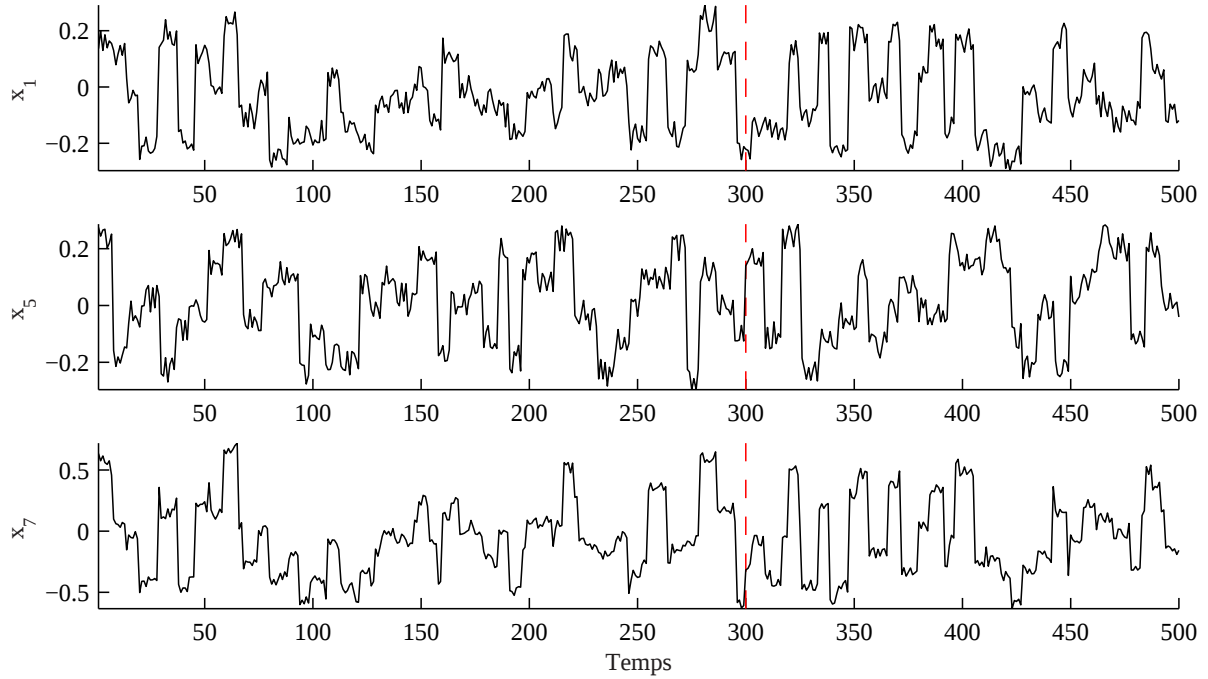
Les mesures simulées des variables  $x_1$ ,  $x_4$  et  $x_7$  sont données, à titre indicatif, sur la figure (fig 2.2).

La matrice de corrélation des variables est donnée par :

$$\Sigma = \begin{pmatrix} 1 & 0.96 & 0.95 & 0.00 & 0.00 & 0.80 & 0.86 \\ 0.96 & 1 & 0.95 & 0.00 & 0.00 & 0.80 & 0.86 \\ 0.95 & 0.95 & 1 & 0.00 & 0.00 & 0.79 & 0.85 \\ 0.00 & 0.00 & 0.00 & 1 & 0.96 & 0.56 & 0.45 \\ 0.00 & 0.00 & 0.00 & 0.96 & 1 & 0.55 & 0.45 \\ 0.80 & 0.80 & 0.79 & 0.56 & 0.55 & 1 & 0.98 \\ 0.86 & 0.86 & 0.85 & 0.45 & 0.45 & 0.98 & 1 \end{pmatrix}$$

Les matrices des valeurs et vecteurs propres sont données par :

$$\Lambda = \begin{pmatrix} 4.69 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 2.15 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0.04 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0.04 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0.03 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0.02 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0.00 \end{pmatrix}$$

Figure 2.2 – Mesures simulées de  $x_1$ ,  $x_4$  et  $x_7$  du premier exemple d'illustration

et

$$P = \begin{pmatrix} 0.42 & -0.24 & 0.41 & -0.36 & 0.59 & -0.31 & 0.02 \\ 0.42 & -0.24 & 0.24 & 0.70 & -0.30 & -0.32 & -0.01 \\ 0.42 & -0.24 & -0.67 & -0.36 & -0.30 & -0.27 & 0.02 \\ 0.16 & 0.62 & 0.38 & -0.35 & -0.48 & -0.26 & 0.06 \\ 0.16 & 0.62 & -0.40 & 0.33 & 0.47 & -0.27 & 0.05 \\ 0.44 & 0.15 & 0.00 & 0.00 & 0.00 & 0.45 & -0.75 \\ 0.45 & 0.07 & 0.01 & 0.03 & 0.00 & 0.60 & 0.65 \end{pmatrix}$$

Les composantes principales sont données par :

$$\begin{aligned} t_1 &= 0.42x_1 + 0.42x_2 + 0.42x_3 + 0.16x_4 + 0.16x_5 + 0.44x_6 + 0.45x_7 \\ t_2 &= -0.24x_1 - 0.24x_2 - 0.24x_3 + 0.62x_4 + 0.62x_5 + 0.15x_6 + 0.07x_7 \\ t_3 &= 0.41x_1 + 0.24x_2 - 0.67x_3 + 0.38x_4 - 0.40x_5 + 0.00x_6 + 0.01x_7 \\ t_4 &= -0.36x_1 + 0.70x_2 - 0.36x_3 - 0.35x_4 + 0.33x_5 + 0.00x_6 + 0.03x_7 \\ t_5 &= 0.59x_1 - 0.30x_2 - 0.30x_3 + 0.48x_4 + 0.47x_5 + 0.00x_6 + 0.00x_7 \\ t_6 &= -0.31x_1 - 0.32x_2 - 0.27x_3 - 0.26x_4 - 0.27x_5 + 0.45x_6 + 0.60x_7 \\ t_7 &= 0.02x_1 - 0.01x_2 + 0.02x_3 + 0.06x_4 + 0.05x_5 - 0.75x_6 + 0.65x_7 \end{aligned}$$

Ainsi, on peut tracer l'évolution des composantes principales  $t_1, \dots, t_7$  de cet exemple.

A partir de la figure (fig 2.3), on remarque que les composantes  $t_3, \dots, t_7$  ne représentent que du bruit alors que les deux premières composantes sont porteuses d'information et sont corrélées avec les variables originelles (car elles sont obtenues par combinaison linéaire de ces dernières).

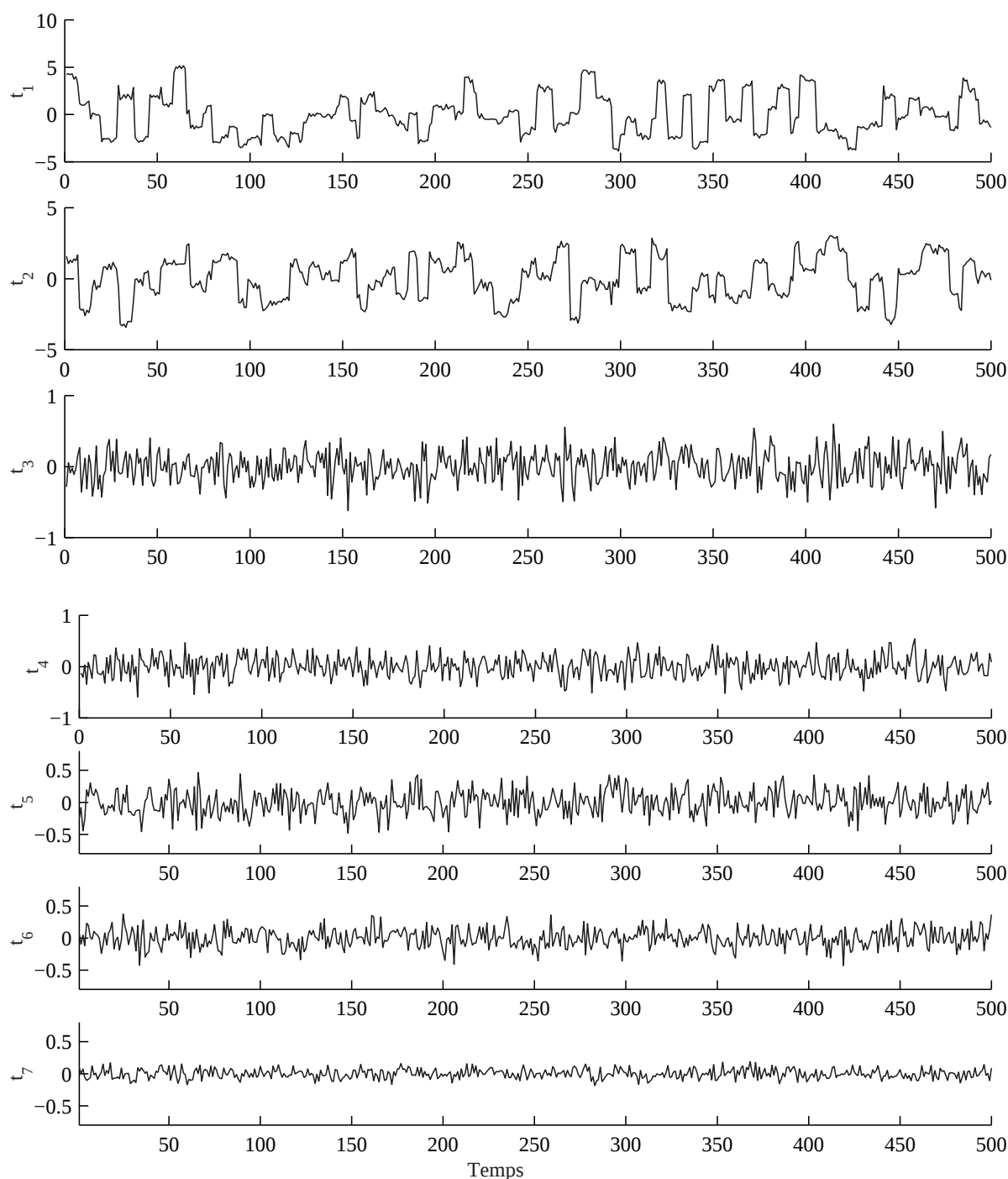


Figure 2.3 – Evolution des composantes principales de l'exemple 1

Cependant, pour l'estimation des variables originales on ne doit conserver que les composantes porteuses d'information significative permettant d'expliquer les différentes variables. Pour cette raison, la suite de ce chapitre sera consacrée au problème de la détermination de la structure du modèle ACP, c'est-à-dire, la détermination du nombre de composantes à conserver ou à retenir dans le modèle.

### 2.3.2 Détermination de la structure du modèle

L'analyse en composantes principales recherche une approximation de la matrice initiale des données  $X$  par une matrice de rang inférieur issue d'une décomposition en valeurs singulières. La question qui se pose alors, et qui a été largement débattue dans la littérature, concerne le choix du nombre de composantes principales qui doit être retenu. De nombreuses règles sont proposées dans la littérature pour déterminer le nombre de composantes à retenir [3, 40, 19, 99]. La plupart de ces méthodes sont heuristiques et donnent un nombre de composantes subjectif. Toutefois, dans le cadre de l'application de l'ACP au diagnostic, le nombre de composantes a un impact significatif sur chaque étape de la procédure de détection et de localisation. Si peu de composantes sont utilisées, on risque de perdre des informations contenues dans les données de départ en projetant certaines variables dans le sous-espace des résidus et donc avoir des erreurs de modélisation qui entachent les résidus ce qui provoque des fausses alarmes. Si par contre beaucoup de composantes sont utilisées, il y a le risque d'avoir des composantes retenues (les composantes correspondantes aux valeurs propres les plus faibles parmi celles retenues dans le modèle) qui sont porteuses de bruit ce qui est indésirable. De plus il y a risque de non détection des défauts, si certaines variables sont projetées dans le sous-espace des composantes principales alors qu'elles doivent être projetées dans le sous-espace résiduel.

Récemment, Qin *et al.* [77] ont proposé une technique basée sur la variance de l'erreur de reconstruction des mesures, ce critère permet de prendre en compte la notion de redondance entre les variables.

Dans la suite nous allons présenter quelques critères utilisés pour la détermination du nombre de composantes et ainsi définir le modèle ACP.

La détermination du nombre de composantes du modèle ACP n'est en fait qu'une procédure de recherche de la structure du modèle. On retrouve dans la littérature plusieurs critères utilisant différentes approches.

#### 2.3.2.1 Critères Heuristiques

Dans cette catégorie de critères on peut citer deux tests.

#### Pourcentage cumulé de la variance totale (PCV)

A la base de cette méthode, notons que chaque composante principale est représentative d'une portion de la variance des mesures du processus étudié. Les valeurs propres de la matrice de corrélation sont les mesures de cette variance et peuvent donc être utilisées dans la sélection du nombre de composantes principales. Pour le choix de  $\ell$ , il faut choisir le pourcentage de la variance totale qu'on veut conserver. Le nombre de composantes est alors le plus petit nombre pris de telle sorte que ce pourcentage soit atteint ou dépassé; les composantes sont choisies successivement dans l'ordre des variances décroissantes. Le pourcentage de variance expliquée par les  $\ell$  premières composantes est donné par :

$$PCV(\ell) = 100 \left( \frac{\sum_{j=1}^{\ell} \lambda_j}{\sum_{j=1}^m \lambda_j} \right) \% \quad (2.36)$$

La variance du bruit étant inconnue a priori, la décision basée seulement sur le pourcentage de la variance expliquée est un peu arbitraire. Sa capacité à fournir le nombre correct de composantes principales dépendra fortement du rapport signal sur bruit.

### Moyenne des valeurs propres

Si les observations constituent un échantillon aléatoire d'individus prélevés dans une population normale à  $m$  dimensions, on peut tester l'égalité des  $(m - \ell)$  dernières valeurs propres. Si cette hypothèse est acceptée, on conserve les  $\ell$  premiers axes et on néglige les  $(m - \ell)$  derniers axes. Cependant l'utilisation de ce test conduit souvent à considérer un nombre élevé de composantes, dont certaines risquent de ne présenter aucun intérêt pratique.

Des règles empiriques peuvent également guider l'utilisateur. Une de ces règles consiste à ne prendre en considération que les composantes pour lesquelles la valeur propre est supérieure à la moyenne arithmétique de toutes les valeurs propres. En particulier, si on travaille sur les données centrées réduites, cela revient à négliger les composantes dont la variance est inférieure à l'unité ( $\frac{1}{m}\text{trace}(\Sigma) = 1$ ). Dans le cas du modèle ACP calculé à partir de la matrice de covariance  $\Sigma$ , la moyenne arithmétique des valeurs propres est donnée par  $\frac{1}{m}\text{trace}(\Sigma)$ .

#### 2.3.2.2 Critères de généralisation

Wax et Kailath [103] ont proposé de déterminer le nombre de composantes principales en utilisant les critères AIC [2] et MDL [81]. Cependant, dans le cadre de l'analyse en composantes principales ces deux critères ont une utilisation très limitée car ils ne sont pas définis pour des valeurs propres nulles ou des valeurs propres de petites amplitudes. Pour plus de détails concernant ces deux critères le lecteur peut consulter l'article de Valle *et al.* [99].

#### 2.3.2.3 Critère de validation croisée

La validation croisée est un critère très populaire pour le choix du nombre de composantes dans un modèle ACP. La base de cette méthode est d'estimer les mesures d'un jeu de données de validation à partir d'un modèle qui a été calculé à partir d'un jeu de données d'identification et de comparer ces estimations avec les valeurs mesurées.

Cette procédure de validation croisée est basée sur la minimisation de la quantité *PRESS*. Cette quantité se présente comme la somme des carrés des erreurs entre les données observées et celles prédites ou estimées par le modèle obtenu à partir d'un jeu d'identification différent.

$$PRESS(\ell) = \frac{1}{Nm} \sum_{k=1}^N \sum_{i=1}^m (\hat{x}_i^{(\ell)}(k) - x_i(k))^2 \quad (2.37)$$

$N$  étant la taille du jeu de validation.

Une version simplifiée de l'algorithme permettant le calcul du nombre de composantes principales par la validation croisée est la suivante :

1. diviser les données en un jeu d'identification et un jeu de validation,
2. réaliser une ACP avec  $\ell$  composantes ( $\ell = 1, \dots, m$ ) sur le jeu d'identification et calculer les critères correspondant sur le jeu de validation  $PRESS(1), \dots, PRESS(m)$ ,
3. la  $\ell$ ème composante pour laquelle le minimum de *PRESS* apparaît sera la dernière composante à retenir et  $\ell$  sera le nombre de composantes principales retenu.

Cependant, Besse *et al.* [3] ont montré théoriquement que, malgré un coût de calcul important, l'usage de la validation croisée en ACP n'apporte pas une règle de décision plus objective que les techniques usuelles heuristiques. Cette équivalence est obtenue par un développement de Taylor du critère de validation croisée en utilisant la théorie des perturbations [3].

### 2.3.2.4 Nombre de composantes pour une meilleure reconstruction

Un critère de sélection du nombre de composantes à retenir dans le modèle ACP a été proposé par Dunia *et al.* [14]. Ce critère cherche le nombre de composantes qui permet de minimiser la variance de l'erreur de reconstruction et que l'on appelle parfois la variance non reconstruite. Ce critère utilisant certaines notions que nous n'avons pas encore définies, il est préférable, dans un premier temps, de présenter le principe de reconstruction. A partir des variables reconstruites, nous allons présenter l'expression de la variance non reconstruite et enfin nous présenterons le critère à minimiser pour la sélection du nombre  $\ell$  de composantes.

#### Principe de reconstruction

Le principe de reconstruction consiste à estimer une des variables du vecteur  $\mathbf{x}(k)$  à un instant donné, notée  $x_i(k)$ , en utilisant toutes les autres variables  $x_j(k)$  au même instant à partir du modèle ACP déjà obtenu. Il existe trois approches différentes de reconstruction qui aboutissent exactement à la même solution.

#### Approche itérative de reconstruction

Une fois l'estimation  $\hat{x}_i$  de la  $i^{eme}$  variable calculée à partir de  $\mathbf{x}$  en utilisant l'équation (2.26), on remplace la  $i^{eme}$  variable  $x_i$  par  $\hat{x}_i$  dans le vecteur des mesures  $\mathbf{x}$ , et on ré-estime cette  $i^{eme}$  variable.

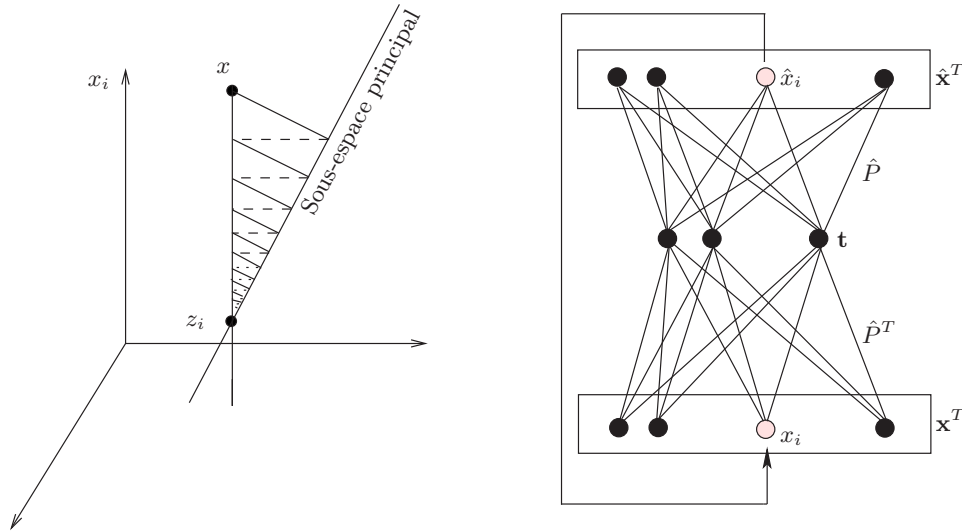


Figure 2.4 – Principe de la méthode itérative de reconstruction

Cette opération est répétée jusqu'à convergence vers la valeur  $z_i$  (fig 2.4). Chaque itération à travers le modèle ACP est une projection orthogonale dans le sous-espace des composantes principales. Ces itérations peuvent être calculées par l'expression suivante :

$$z_i^{(iter)} = [\mathbf{x}_{-i}^T \ z_i^{(iter-1)} \ \mathbf{x}_{+i}^T] \mathbf{c}_i = c_{ii} z_i^{(iter-1)} + [c_{-i}^T \ 0 \ c_{+i}^T] \mathbf{x} \text{ avec } z_i^{(0)} = x_i \quad (2.38)$$

où  $\hat{\mathbf{C}} = \hat{\mathbf{P}}\hat{\mathbf{P}}^T = [c_1 \ c_2 \ \dots \ c_m]$  et  $z_i$  est la valeur reconstruite, à convergence, de la mesure de la  $i^{eme}$  variable.  $z_i^{(iter-1)}$  est la valeur de la mesure estimée par l'ACP à l'itération précédente

et  $z_i^{(iter)}$  est la nouvelle valeur de la mesure estimée par l'ACP à partir du vecteur  $\hat{\mathbf{x}}_i^{(iter-1)}$  qui représente le vecteur  $\mathbf{x}$  dont la  $i^{eme}$  composante a été reconstruite à l'itération précédente.

$$c_i^T = [c_{1i} \ c_{2i} \ \dots \ c_{mi}] = [c_{-i}^T \ c_{ii} \ c_{+i}^T] \quad (2.39)$$

et  $\mathbf{x}^T$  est un vecteur ligne contenant les mesures des capteurs à un instant donné et les indices  $+i$  et  $-i$  désignent les vecteurs formés par les  $(i-1)$  premiers et les  $(m-i)$  derniers éléments du vecteur originel, respectivement. A la convergence de l'algorithme (2.38), la solution  $z_i$  vérifie :

$$z_i(1 - c_{ii}) = [c_{-i}^T \ 0 \ c_{+i}^T] \mathbf{x} \quad (2.40)$$

Ainsi,  $z_i$  sera donnée par :

$$z_i = \frac{[c_{-i}^T \ 0 \ c_{+i}^T] \mathbf{x}}{1 - c_{ii}} \quad (c_{ii} \neq 1) \quad (2.41)$$

et  $z_i = x_i$  pour  $c_{ii} = 1$ . Dans ce cas :

$$c_i = [0 \dots 1 \dots 0]^T \quad (2.42)$$

est la  $i^{eme}$  colonne de la matrice identité. Ceci correspond au cas où  $x_i$  ne serait pas corrélée avec les autres variables et la variable ne peut donc pas être reconstruite à partir des autres.

### Approche des valeurs manquantes

Si on traite l'échantillon délivré par un capteur défaillant comme une valeur manquante, l'approche de remplacement des données manquantes peut être utilisée dans la reconstitution des mesures délivrées par ce capteur. La méthode de remplacement des valeurs manquantes en utilisant l'ACP est la suivante. Si le  $i^{eme}$  capteur est défaillant, la mesure est retirée du vecteur des échantillons. Le vecteur des échantillons modifié sera :

$$\mathbf{x}^{(i)} = [x_{-i}^T \ x_{+i}^T]^T \in \mathbb{R}^{m-1} \quad (2.43)$$

Le vecteur des composantes principales peut être obtenu à partir de l'équation suivante :

$$\hat{\mathbf{t}} = \arg \min_{\mathbf{t}} \|\mathbf{x}^{(i)} - \hat{P}_i \mathbf{t}\|^2 \quad (2.44)$$

sachant que  $\hat{\mathbf{x}}^{(i)} = \hat{P}_i \mathbf{t}$ , c'est le vecteur des composantes principales qui minimise le  $SPE$ .

$$\hat{\mathbf{t}} = (\hat{P}_i^T \hat{P}_i)^{-1} \hat{P}_i^T \mathbf{x}^{(i)} \quad (2.45)$$

où  $\hat{P}_i = \begin{bmatrix} \hat{P}_{-i} \\ \hat{P}_{+i} \end{bmatrix} \in \mathbb{R}^{(m-1) \times \ell}$

La valeur reconstruite de  $\mathbf{x}^{(i)}$  est :

$$z_i = \xi_i^T \hat{P} \hat{\mathbf{t}} = \xi_i^T \hat{P} (\hat{P}_i^T \hat{P}_i)^{-1} \hat{P}_i^T \mathbf{x}^{(i)} \quad (2.46)$$

La reconstruction, par l'approche itérative, donnée par l'équation (2.41) est équivalente à celle de l'équation (2.46). La seule différence est que l'approche itérative ne nécessite pas d'inversion matricielle [14].

### Approche par optimisation du $SPE$

Cette approche a été présentée par Wise [112]. La procédure d'optimisation cherche à trouver la valeur reconstruite qui minimise le  $SPE$  et donc l'erreur d'estimation :

$$z_i = \arg \min_{z_i} \left\| (I - \hat{C}) \begin{bmatrix} x_{-i}^T & z_i & x_{+i}^T \end{bmatrix}^T \right\|^2 \quad (2.47)$$

En posant la dérivée du  $SPE$  par rapport à  $z_i$  égale à zéro, on retrouve la condition nécessaire pour avoir un extremum :

$$\xi_i^T (I - C) \mathbf{x}_i = \begin{bmatrix} x_{-i}^T & z_i & x_{+i}^T \end{bmatrix} (\xi_i - c_i) = 0 \quad (2.48)$$

où  $\mathbf{x}_i = [x_1 \dots x_{i-1} z_i x_{i+1} \dots x_m]^T$  représente le vecteur  $\mathbf{x}$  dont la  $i^{eme}$  variable a été reconstruite. Cette équation conduit à la même solution que l'équation (2.41) :

$$\begin{bmatrix} x_{-i}^T & z_i & x_{+i}^T \end{bmatrix} \left( [0 \dots 1 \dots 0]^T - [c_{i1} \dots c_{ii} \dots c_{im}]^T \right) = 0 \quad (2.49)$$

$$z_i (1 - c_{ii}) - \begin{bmatrix} c_{-i}^T & 0 & c_{+i}^T \end{bmatrix} \mathbf{x}_i = z_i (1 - c_{ii}) - \begin{bmatrix} c_{-i}^T & 0 & c_{+i}^T \end{bmatrix} \mathbf{x} = 0 \quad (2.50)$$

Toutefois, il faut noter que la condition nécessaire et suffisante de reconstruction, à partir de l'équation (2.41), est que  $c_{ii} \neq 1$ , ce qui peut se traduire par  $\tilde{\xi}_i \neq 0$ , où  $\tilde{\xi}_i = (I - \hat{C})\xi_i$ .

Ainsi, nous avons présenté les trois approches de reconstruction, mais il reste un certain détail à expliciter. Il faut noter ici que l'erreur d'estimation de toute valeur reconstruite est strictement nulle. Pour démontrer cette propriété, nous allons considérer le cas où  $z_i$  est donnée par l'équation (2.41), ainsi  $\hat{z}_i$  sera donnée par l'expression :

$$\hat{z}_i = \begin{bmatrix} c_{-i}^T & c_{ii} & c_{+i}^T \end{bmatrix} \mathbf{x}_i \quad (2.51)$$

Partant de l'équation (2.41) nous pouvons écrire que :

$$z_i = \begin{bmatrix} c_{-i}^T & 0 & c_{+i}^T \end{bmatrix} \mathbf{x} + c_{ii} z_i \quad (2.52)$$

En comparant les deux expressions des équations (2.52) et (2.51) on constate qu'elles sont identiques et ainsi par reconstruction l'erreur d'estimation de la variable reconstruite est nulle.

$$e_i = z_i - \hat{z}_i = 0 \quad (2.53)$$

Cette propriété de la reconstruction est très intéressante et a été à la base d'une méthode de localisation de défauts proposée par Dunia *et al.* [14].

**Remarque 2.3.1** Nous avons présenté le principe de reconstruction dans le cas d'une seule variable, mais il peut être généralisé à la reconstruction de plusieurs variables.

### Principe de la variance non reconstruite

La notion de variance non reconstruite est basée sur l'estimation de l'information délivrée par le capteur  $i$  en utilisant toutes les autres mesures des capteurs. Il y a toujours une partie de la variation des mesures qui ne peut être reconstruite en utilisant les autres capteurs ; cette partie est l'erreur de reconstruction. L'erreur de reconstruction de la  $i^{eme}$  variable est donnée par :

$$\xi_i^T (\mathbf{x} - \mathbf{x}_i) = \left( \xi_i^T - \frac{\begin{bmatrix} c_{-i}^T & 0 & c_{+i}^T \end{bmatrix}}{1 - c_{ii}} \right) \mathbf{x} = \frac{\tilde{\xi}_i^T \mathbf{x}}{\tilde{\xi}_i^T \tilde{\xi}_i} \quad (2.54)$$

où  $\tilde{\xi}_i = (I - \hat{C})\xi_i$ , et  $\tilde{\xi}_i^T \tilde{\xi}_i = (1 - c_{ii})$ .

La norme de l'erreur de reconstruction de la  $i^{eme}$  variable est donnée par :

$$\|\mathbf{x} - \mathbf{x}_i\| = \frac{|\tilde{\xi}_i^T \mathbf{x}|}{\tilde{\xi}_i^T \tilde{\xi}_i} \quad (2.55)$$

A partir des équations (2.54) et (2.55), nous observons que :

1.  $\mathcal{E}(\mathbf{x} - \mathbf{x}_i) = 0$  puisque  $\mathcal{E}(\mathbf{x}) = 0$ . En d'autres termes,  $\mathbf{x}_i$  est une estimation non biaisée de  $\mathbf{x}$ ,
2. la variance de  $(\mathbf{x} - \mathbf{x}_i)$  est non nulle dans la direction de reconstruction  $\xi_i$ ,
3. l'amplitude de l'erreur de reconstruction  $(\mathbf{x} - \mathbf{x}_i)$  dépend du nombre de composantes principales retenues dans le modèle ACP. En effet,  $\tilde{\xi}_i$  dépend de  $\tilde{C} = (I - \hat{C})$  et donc de  $\ell$  (car  $(I - \hat{C}) = (I - \hat{P}\hat{P}^T)$ ).

Les deux derniers points indiquent qu'on peut déterminer le nombre de composantes principales à retenir par minimisation de l'erreur de reconstruction. En d'autre termes, on peut trouver le nombre de composantes permettant de réaliser la meilleure reconstruction. La variance non reconstruite dans la direction  $\xi_i$  (reconstruction de la  $i^{eme}$  variable) est notée  $\rho_i$  et représente la variance de la projection de  $(\mathbf{x} - \mathbf{x}_i)$  dans la direction  $\xi_i$ .

$$\rho_i = var\{\xi_i^T(\mathbf{x} - \mathbf{x}_i)\} = \mathcal{E}\{(\mathbf{x} - \mathbf{x}_i)(\mathbf{x} - \mathbf{x}_i)^T\} = \frac{\tilde{\xi}_i^T \mathcal{E}\{\mathbf{x}\mathbf{x}^T\} \tilde{\xi}_i}{(\tilde{\xi}_i^T \tilde{\xi}_i)^2} \quad (2.56)$$

$\mathcal{E}\{\mathbf{x}\mathbf{x}^T\} = \Sigma$  est la matrice de corrélation des données.

En tenant compte du fait que  $\hat{C}\Sigma = \Sigma\hat{C}$  et que  $\hat{C}\Sigma = \hat{C}\Sigma\hat{C}$ ,  $\Sigma$  peut s'écrire sous la forme :

$$\Sigma = \hat{C}\Sigma\hat{C} + (I - \hat{C})\Sigma(I - \hat{C}) = \hat{\Sigma} + \tilde{\Sigma} \quad (2.57)$$

où  $\hat{\Sigma} = \mathcal{E}\{\hat{\mathbf{x}}\hat{\mathbf{x}}^T\}$  et  $\tilde{\Sigma} = \mathcal{E}\{\tilde{\mathbf{x}}\tilde{\mathbf{x}}^T\}$  sont respectivement les parties modélisées et non modélisées de  $\Sigma$ . En remplaçant l'expression  $\mathcal{E}\{\mathbf{x}\mathbf{x}^T\}$  de l'équation (2.56), nous obtenons :

$$\rho_i = \frac{\tilde{\xi}_i^T \Sigma \tilde{\xi}_i}{(\tilde{\xi}_i^T \tilde{\xi}_i)^2} = \frac{\tilde{\xi}_i^T \tilde{\Sigma} \tilde{\xi}_i}{(\tilde{\xi}_i^T \tilde{\xi}_i)^2} \quad (2.58)$$

La variance de l'erreur de reconstruction peut être utilisée pour la détermination du nombre de composantes à retenir dans le modèle ACP. Ceci est effectué par minimisation de cette variance par rapport au nombre de composantes et pour toutes les variables. Pour illustrer cette minimisation, l'identité suivante est utilisée :

$$\|\hat{\xi}_i\|^2 + \|\tilde{\xi}_i\|^2 = \|\xi_i\|^2 = 1 \quad (2.59)$$

où  $\hat{\xi}_i = \hat{C}\xi_i$  est la projection de  $\xi_i$  dans le sous-espace des composantes principales, et  $\tilde{\xi}_i = (I - \hat{C})\xi_i$  est la projection de  $\xi_i$  dans le sous-espace des résidus. En multipliant (2.59) par  $\rho_i$ , on obtient :

$$\rho_i = \rho_i \|\hat{\xi}_i\|^2 + \rho_i \|\tilde{\xi}_i\|^2 \quad (2.60)$$

où  $\rho_i \|\hat{\xi}_i\|^2$  est la variance non reconstruite dans le sous-espace des composantes principales et sera notée  $\hat{\rho}_i$ .  $\rho_i \|\tilde{\xi}_i\|^2$  est la variance non reconstruite dans le sous-espace des résidus et sera noté par  $\tilde{\rho}_i$ . L'équation (2.60) devient alors :

$$\rho_i = \hat{\rho}_i + \tilde{\rho}_i \quad (2.61)$$

En comparant les deux équations (2.60) et (2.61), nous obtenons :

$$\tilde{\rho}_i = \rho_i \left\| \tilde{\xi}_i \right\|^2 = \tilde{\xi}_i^{\circ T} \tilde{\Sigma} \tilde{\xi}_i^{\circ} \quad (2.62)$$

où  $\tilde{\xi}_i^{\circ} = \tilde{\xi}_i / \left\| \tilde{\xi}_i \right\|$  et

$$\hat{\rho}_i = \tilde{\rho}_i \frac{\left\| \hat{\xi}_i \right\|^2}{\left\| \tilde{\xi}_i \right\|^2} \quad (2.63)$$

Dunia *et al.* [17] ont montré que  $\tilde{\rho}_i$  est monotone décroissante avec  $\ell$  et  $\hat{\rho}_i$  tend vers l'infini pour  $\ell = m$  ;  $\rho_i$  doit nécessairement avoir un minimum dans l'intervalle  $[1, m]$ . Il est possible que le minimum soit pour  $\ell = 1$  ; toutefois, pour une meilleure reconstruction,  $\ell$  doit être inférieur à  $m$ . La réduction de  $\rho_i$ , en choisissant  $\ell$ , améliore la reconstruction. Donc, un problème d'optimisation peut être formulé où l'objectif est de minimiser  $\rho_i$  par rapport au nombre de composantes principales  $\ell$  :

$$\min_{\ell} \rho_i(\ell) \quad (2.64)$$

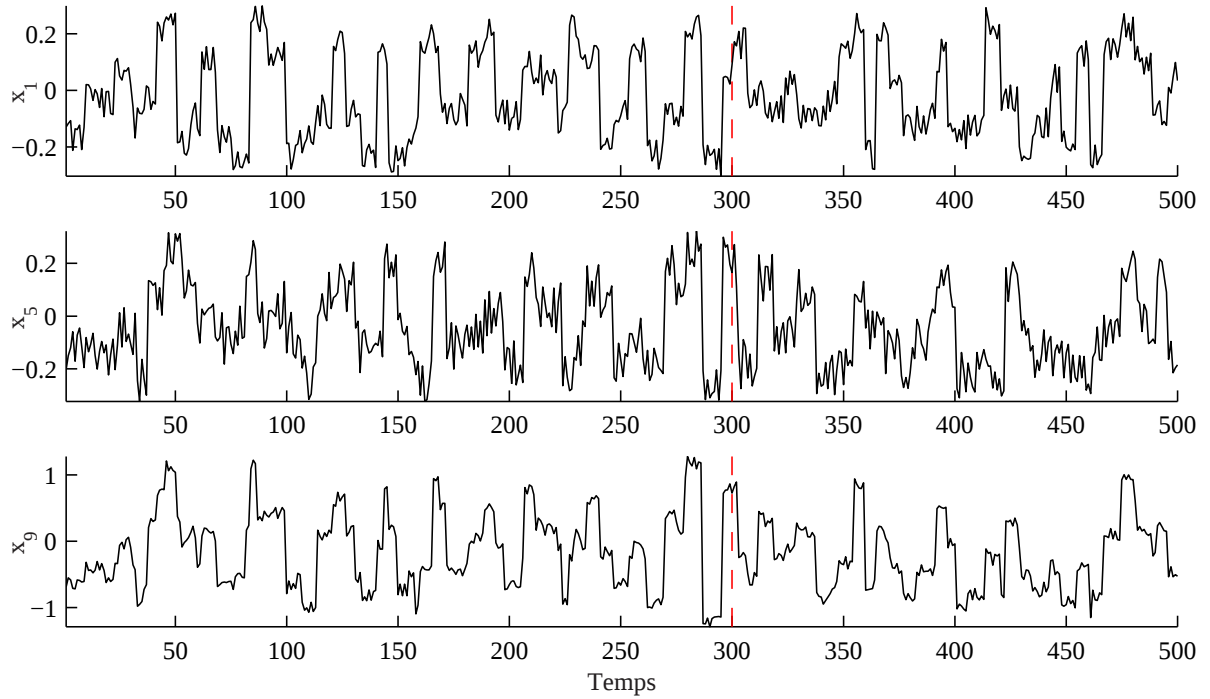
Puisque cette équation fournit la valeur de  $\ell$  optimale seulement pour la reconstruction de la  $i^{eme}$  variable, il est nécessaire de définir un critère considérant toutes les variables  $x_i$  ( $i = 1, \dots, m$ ) :

$$\min_{\ell} \sum_{i=1}^m \rho_i(\ell) \quad (2.65)$$

Dans ce deuxième exemple, nous allons essayer de nous rapprocher de la réalité (cas pratiques) en introduisant des non linéarités douces dans l'exemple 1. Nous disposons de dix variables qui sont décrites par le système d'équations suivant :

$$\begin{aligned} x_1 &= u_1 + \varepsilon_1 \\ x_2 &= u_1 + \varepsilon_2 \\ x_3 &= u_1 + \varepsilon_3 \\ x_4 &= u_2 + \varepsilon_4 \\ x_5 &= u_2 + \varepsilon_5 \\ x_6 &= u_2^2 + \varepsilon_6 \\ x_7 &= u_1 + u_2^2 + \varepsilon_7 \\ x_8 &= 5u_1^2 + 3u_2^2 + \varepsilon_8 \\ x_9 &= 2u_1 + 3u_2 + \varepsilon_9 \\ x_{10} &= u_1^2 + u_2 + \varepsilon_{10} \end{aligned} \quad (2.66)$$

Dans l'ensemble des variables simulées, nous avons une redondance directe d'ordre 3 entre les variables  $x_1$ ,  $x_2$  et  $x_3$  et une autre redondance directe d'ordre 2 entre  $x_4$  et  $x_5$ . Nous avons également des redondances analytiques linéaires et non linéaires. Les non linéarités ne sont pas très fortes, comme on peut le constater pour les variables  $x_6$  et  $x_8$  car les variables sur lesquelles les non linéarités ont été introduites ont une plage de variation comprise entre  $-0.25$  et  $+0.25$ . Nous disposons de 500 échantillons pour chaque variable, les 300 premiers échantillons ont été

Figure 2.5 – Evolution des variables  $x_1$ ,  $x_5$  et  $x_9$  de l'exemple 2

utilisés pour l'identification du modèle ACP et le reste des échantillons a été utilisé pour le diagnostic (fig 2.5).

Les figures (fig 2.6) et (fig 2.7) présentent les résultats de la comparaison entre quatre indices de sélection du nombre de composantes principales appliqués sur les deux exemples 1 et 2. Le critère basé sur la variance de l'erreur de reconstruction (variance non reconstruite) est le critère le plus intéressant pour les objectifs de diagnostic car il tient compte des redondances qui existent entre les différentes variables en utilisant le principe de reconstruction.

Pour les deux exemples de simulation que nous avons présentés, le critère des valeurs propres donne le même nombre de composantes que le critère VNR (Variance Non Reconstruite). Cependant, pour des applications réelles, ce critère a tendance à souestimer le nombre exact de composantes à retenir dans le modèle ACP [99].

### 2.3.3 Sélection des capteurs pour une meilleure reconstruction

Nous avons décrit  $\rho_i$  comme étant un paramètre qui quantifie la quantité de variance qui n'est pas capturée par  $z_i$  en estimant  $x_i$ . Toutefois,  $\rho_i$  peut être très grand, après minimisation, pour des capteurs particuliers. Un critère plus restrictif que la condition de reconstruction est nécessaire pour déterminer si  $z_i$  est une estimation adéquate de  $x_i$ . Pour cela nous définissons une valeur de référence pour déterminer si une valeur  $\rho_i$  particulière donne une reconstruction acceptable de la  $i^{eme}$  variable. Cette référence est obtenue en utilisant la moyenne  $\bar{\mathbf{x}}$  du vecteur des échantillons à la place du vecteur reconstruit  $\mathbf{x}_i$  :

$$\rho_i(0) = var \left\{ \xi_i^T (\mathbf{x} - \bar{\mathbf{x}}) \right\} = \xi_i^T \Sigma \xi_i \quad (2.67)$$

Si  $\mathbf{x}$  est normalisé, sa moyenne est nulle et sa variance vaut un. Ce cas correspond à l'utilisation

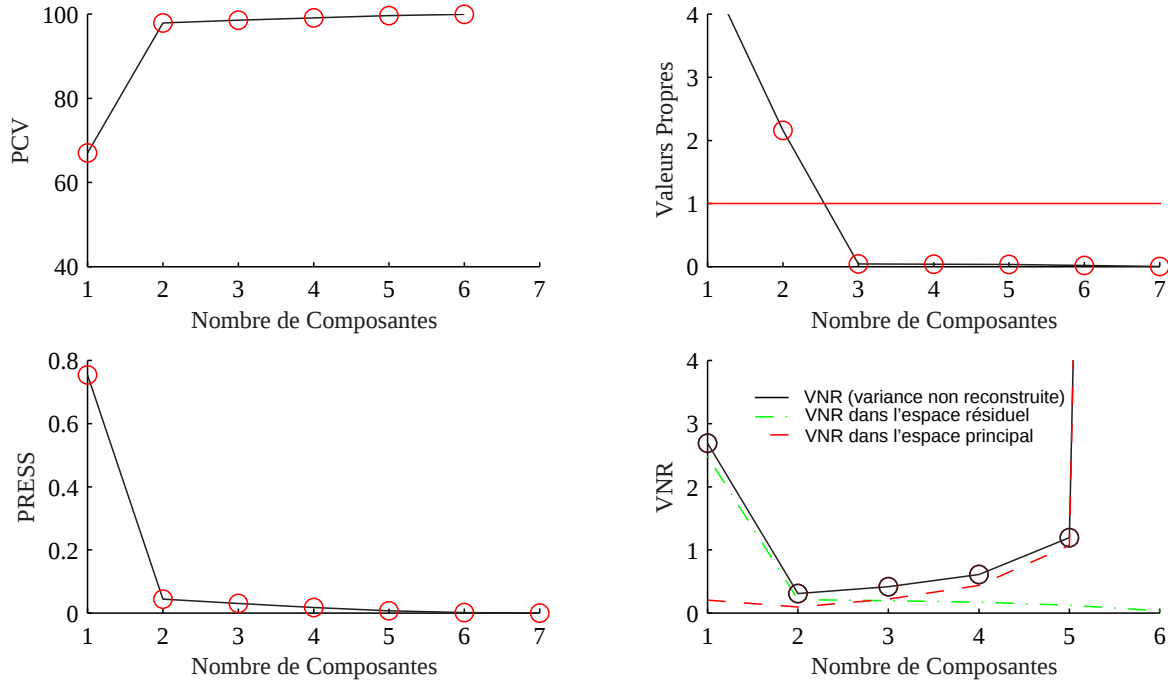


Figure 2.6 – Evolution des différents critères de sélection du nombre de composantes en fonction du nombre de composantes pour l'exemple 1

d'un modèle avec aucune composante principale. En particulier,  $\rho_i(0) = \tilde{\rho}_i(0)$  et  $\hat{\rho}_i = 0$ . Ceci peut s'expliquer par le fait que l'espace résiduel est de dimension  $m$ , ainsi toute la variance est projetée dans le sous-espace résiduel.

En ajoutant des composantes principales, il est possible que  $\rho_i(\ell) \geq \rho_i(0)$  si la croissance de  $\hat{\rho}_i$  est supérieure à la décroissance de  $\tilde{\rho}_i$ . Si la variance non reconstruite  $\rho_i(\ell)$  est supérieure à  $\rho_i(0)$ , la moyenne  $\bar{\mathbf{x}}$  donne une variance non reconstruite pour l'estimation de  $\mathbf{x}_i$  plus faible que celle obtenue avec  $\ell$  composantes. Donc, Le modèle ACP obtenu à partir des données ne donne pas une estimation de  $\mathbf{x}$  meilleure que  $\bar{\mathbf{x}}$ . Dans ce cas, l'utilisation de  $\bar{\mathbf{x}}$  comme valeur reconstruite est la meilleure solution, qui peut être interprétée comme l'utilisation de zéro composante.

Comme la reconstruction d'une variable dépend des autres variables utilisées dans le modèle ACP, la qualité d'une reconstruction est d'autant meilleure que la variable reconstruite est plus corrélée avec les autres variables et la variance de reconstruction est plus petite. Si la  $i$ ème variable n'est pas corrélée avec les autres, la variance  $\rho_i$  tend à croître car la valeur  $\rho_i$  est plus grande que  $\rho_i(0)$ . Qin et al. [77] proposent une procédure pour la détermination des  $m$  variables à retenir pour la surveillance et le nombre  $\ell$  de composantes à retenir en utilisant la variance non reconstruite (fig 2.8).

1. Prendre  $m$  le nombre total de variables comme valeur initiale.
2. Estimer la matrice de covariance  $\Sigma$  à partir des données. Minimiser la fonction  $\sum_{i=1}^m \rho_i$  par rapport à  $\ell$ .
3. Écarter, de l'ensemble des capteurs utilisés pour la surveillance, les variables pour lesquelles les variances non reconstruites individuelles sont supérieures à  $\xi_i^T \Sigma \xi_i$  (qui seront égales à un si les données sont centrées réduites),

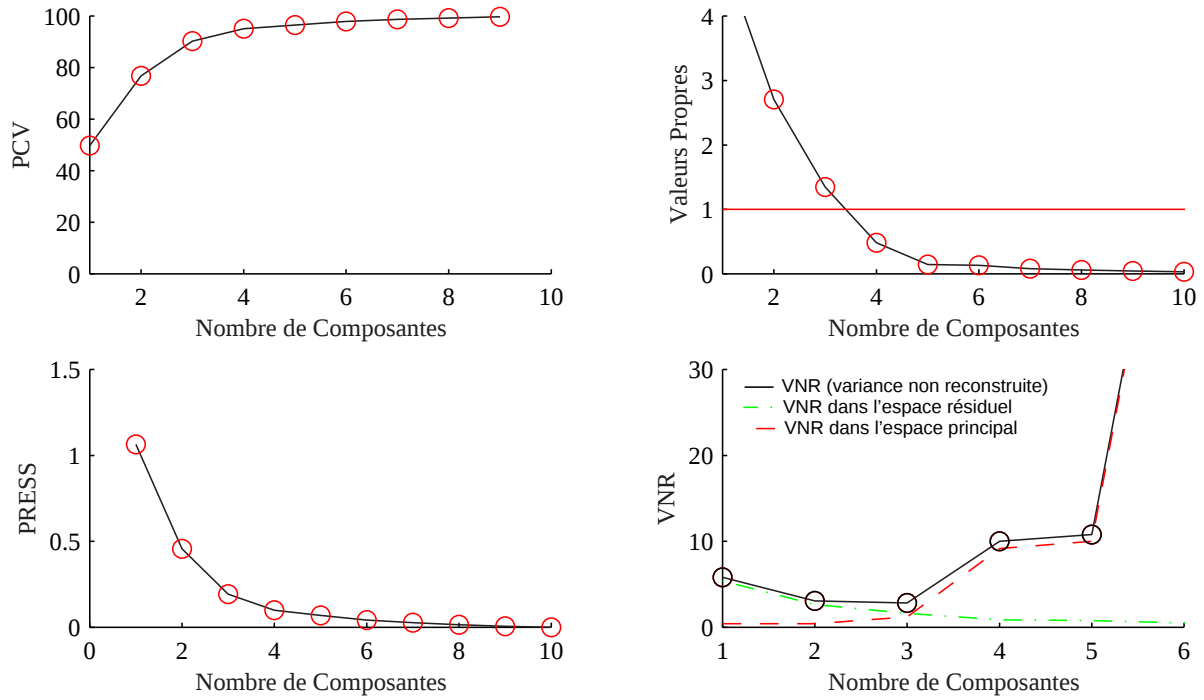


Figure 2.7 – Evolution des différents critères de sélection du nombre de composantes en fonction du nombre de composantes pour l'exemple 2

4. Mettre à jour le nombre des variables utilisées pour la surveillance, si  $m$  reste le même, la procédure d'optimisation est terminée, sinon aller à l'étape 2.

Pour illustrer le principe de cette méthode, considérons l'exemple 2 précédent dans lequel on a rajouté une variable  $x_{11}$  qui est corrélée avec aucune autre variable. Les résultats sont représentés sur le tableau (tab 2.1), où le nombre de composantes qui minimise la fonction coût est  $\ell = 3$ . Pour ce nombre de composantes, on remarque bien que la variance de l'erreur de reconstruction de la variable  $x_{11}$  est supérieure à  $\xi_i^T \Sigma \xi_i = 1$ . Ceci implique que cette variable est corrélée avec aucune autre variable et ainsi elle est écartée de l'ensemble des variables à surveiller et on ne conserve que les dix premières variables. Le tableau (tab 5.3) présente les variances non reconstruites pour l'ensemble des dix variables qui restent. Il montre bien que le minimum est obtenu pour  $\ell = 3$  et que les variances individuelles non reconstruites de toutes les variables calculées en se basant sur un modèle à trois composantes sont toutes inférieures à 1. Ceci implique d'après la procédure de sélection des variables présentée que ces dix variables seront conservées et peuvent être utilisées pour la surveillance de ce processus.

Une fois le nombre de composantes  $\ell$  déterminé, le modèle ACP est ainsi identifié et la matrice de donnée  $X$  peut être approximée à partir des  $\ell$  premières composantes principales correspondant aux  $\ell$  plus grandes valeurs propres de la matrice  $\Sigma$ .

$$\hat{X} = \sum_{i=1}^{\ell} T_i p_i^T = \sum_{i=1}^{\ell} X p_i p_i^T \quad (2.68)$$

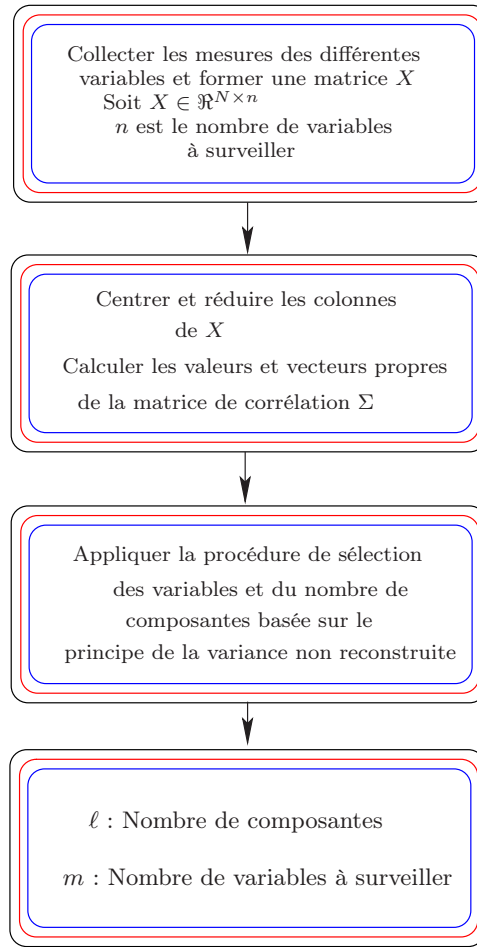


Figure 2.8 – Différentes étapes pour la détermination du nombre de composantes (modèle ACP) et des variables à surveiller.

Sachant que  $\sum_{i=1}^{\ell} p_i p_i^T = \hat{P} \hat{P}^T$  et que l'on notera par  $\hat{C}$ . L'estimation de  $X$  sera donnée par :

$$\hat{X} = X \hat{C} \quad (2.69)$$

Les deux figures (fig 2.9) et (fig 2.10) présentent les mesures et les estimations des trois premières variables de l'exemple 1 et de l'exemple 2, respectivement. Ainsi, les deux modèles identifiés dans le cas des deux exemples traités donnent une estimation assez correcte des mesures.

## 2.4 Conclusion

Dans ce chapitre nous avons présenté le principe de l'analyse en composantes principales linéaires. L'idée de base de l'ACP est de réduire la dimension de la matrice des données, en retenant le plus possible les variations présentes dans le jeu de données de départ. Cette réduction ne sera possible que si les variables initiales ne sont pas indépendantes et ont des coefficients de corrélation entre elles non nuls. Ces variables initiales sont transformées en de nouvelles variables, appelées composantes principales. Elles sont obtenues par combinaisons linéaires des précédentes et sont ordonnées et non corrélées entre elles.

| $i$                         | $\ell$ |      |             |       |       | $\xi_i^T \Sigma \xi_i$ | $\rho_i(3) < \xi_i^T \Sigma \xi_i$ |
|-----------------------------|--------|------|-------------|-------|-------|------------------------|------------------------------------|
|                             | 1      | 2    | 3           | 4     | 5     |                        |                                    |
| 1                           | 0.37   | 0.10 | 0.07        | 0.07  | 0.07  | 1                      | ✓                                  |
| 2                           | 0.39   | 0.14 | 0.11        | 0.11  | 0.10  | 1                      | ✓                                  |
| 3                           | 0.36   | 0.10 | 0.08        | 0.08  | 0.08  | 1                      | ✓                                  |
| 4                           | 0.71   | 0.22 | 0.18        | 0.18  | 0.17  | 1                      | ✓                                  |
| 5                           | 0.67   | 0.22 | 0.19        | 0.19  | 0.18  | 1                      | ✓                                  |
| 6                           | 0.99   | 0.90 | 0.79        | 0.79  | 2.91  | 1                      | ✓                                  |
| 7                           | 0.41   | 0.06 | 0.05        | 0.05  | 0.05  | 1                      | ✓                                  |
| 8                           | 1.00   | 0.95 | 0.70        | 1.01  | 6.10  | 1                      | ✓                                  |
| 9                           | 0.16   | 0.05 | 0.04        | 0.04  | 0.04  | 1                      | ✓                                  |
| 10                          | 0.71   | 0.27 | 0.18        | 0.18  | 0.18  | 1                      | ✓                                  |
| 11                          | 0.99   | 0.99 | <b>1.22</b> | 64.84 | 77.39 | 1                      | ×                                  |
| $\sum_{i=1}^m \rho_i(\ell)$ | 6.76   | 4.00 | 3.61        | 67.54 | 57.27 |                        |                                    |

Table 2.1 – Variances des erreurs de reconstruction des différentes variables avec  $m = 11$  et  $\min_{\ell} \sum_{i=1}^m \rho_i(\ell) = 3$

| $i$                         | $\ell$ |      |      |      |       | $\xi_i^T \Sigma \xi_i$ | $\rho_i(3) < \xi_i^T \Sigma \xi_i$ |
|-----------------------------|--------|------|------|------|-------|------------------------|------------------------------------|
|                             | 1      | 2    | 3    | 4    | 5     |                        |                                    |
| 1                           | 0.37   | 0.10 | 0.07 | 0.07 | 0.07  | 1                      | ✓                                  |
| 2                           | 0.39   | 0.14 | 0.11 | 0.10 | 0.10  | 1                      | ✓                                  |
| 3                           | 0.36   | 0.10 | 0.08 | 0.08 | 0.08  | 1                      | ✓                                  |
| 4                           | 0.71   | 0.22 | 0.18 | 0.18 | 0.35  | 1                      | ✓                                  |
| 5                           | 0.67   | 0.22 | 0.19 | 0.18 | 0.54  | 1                      | ✓                                  |
| 6                           | 0.99   | 0.90 | 0.76 | 2.90 | 2.88  | 1                      | ✓                                  |
| 7                           | 0.41   | 0.06 | 0.05 | 0.05 | 0.05  | 1                      | ✓                                  |
| 8                           | 1.00   | 0.95 | 0.78 | 6.18 | 6.46  | 1                      | ✓                                  |
| 9                           | 0.16   | 0.05 | 0.06 | 0.05 | 0.05  | 1                      | ✓                                  |
| 10                          | 0.71   | 0.27 | 0.20 | 0.18 | 0.18  | 1                      | ✓                                  |
| $\sum_{i=1}^m \rho_i(\ell)$ | 5.77   | 3.01 | 2.48 | 9.97 | 10.76 |                        |                                    |

Table 2.2 – Variances des erreurs de reconstruction des différentes variables avec  $m = 10$  et  $\min_{\ell} \sum_{i=1}^m \rho_i(\ell) = 3$

L'analyse en composantes principales cherche à identifier les vecteurs propres orthonormaux et les valeurs propres de la matrice de dispersion des variables originelles. Les vecteurs propres orthonormaux sont utilisés pour construire les composantes principales, et les valeurs propres sont les variances des composantes principales correspondantes [48].

Pour l'utilisation de l'ACP dans le domaine du diagnostic, un modèle ACP est défini globalement par la matrice des premiers vecteurs propres de la matrice de corrélations des données. Cette matrice permet à la fois de définir la projection permettant d'avoir les composantes principales et la projection inverse permettant d'estimer les données originelles.

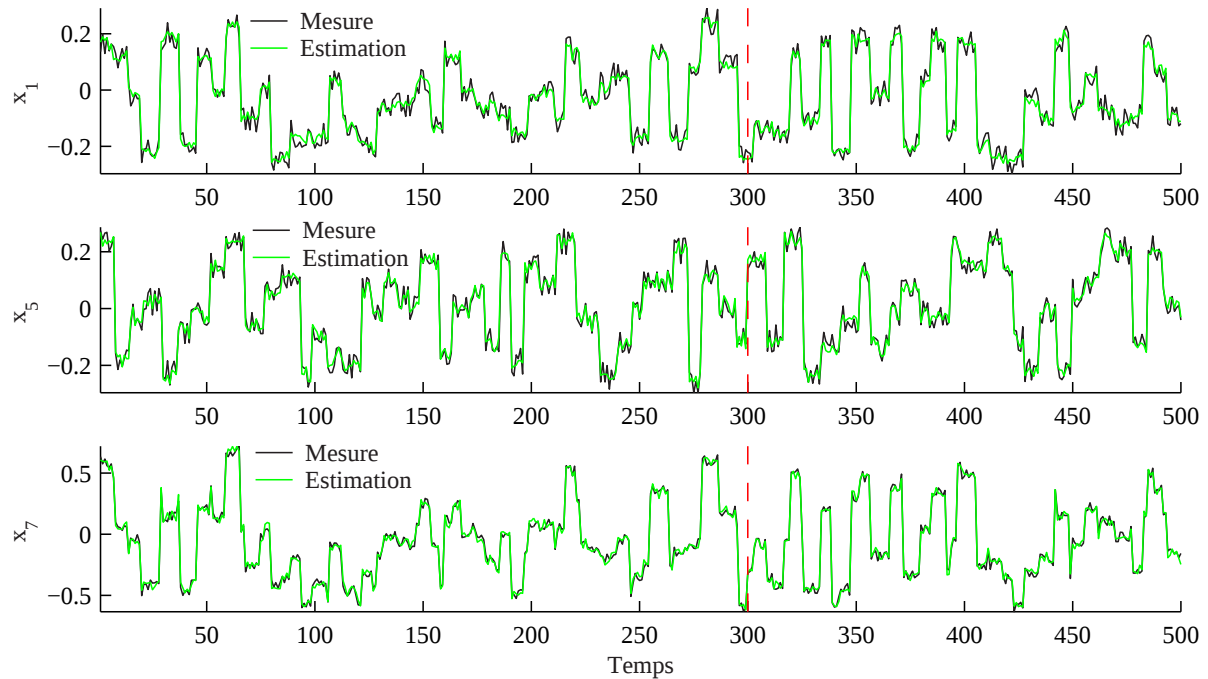


Figure 2.9 – Evolution des mesures et des estimations des trois variables  $x_1$ ,  $x_4$  et  $x_7$  de l'exemple 1

L'identification du modèle ACP nécessite la détermination du nombre de composantes à retenir ainsi que l'estimation des paramètres de ce dernier.

Pour la détermination du nombre de composantes à retenir dans le modèle ACP, plusieurs critères de sélection du nombre de composantes ont été présentés. Le critère de sélection du nombre de composantes basé sur le principe de reconstruction est très intéressant pour des objectifs de diagnostic, car le principe de reconstruction exploite la redondance qui existe entre les différentes variables. Ceci peut s'expliquer par le fait que le modèle obtenu en retenant le nombre de composantes choisi permet de reconstruire toutes les variables à partir des autres. Ainsi, la détermination du nombre de composantes est effectuée par minimisation de la variance de l'erreur de reconstruction.

Avant la procédure de surveillance, une étape concernant la sélection de variables à surveiller est indispensable ; cette étape permet de ne conserver que les variables les plus corrélées et ainsi améliorer la qualité du modèle utilisé pour la surveillance.

Une fois que les variables à surveiller sont sélectionnées et que le nombre de composantes à retenir est déterminé, le modèle ACP est identifié. La procédure de détection et localisation de défaut peut être effectuée par génération des indicateurs de défauts (résidus) en comparant le comportement observé du processus donné par les variables mesurées et le comportement prévu donné par les estimations de ces variables à partir du modèle ACP.

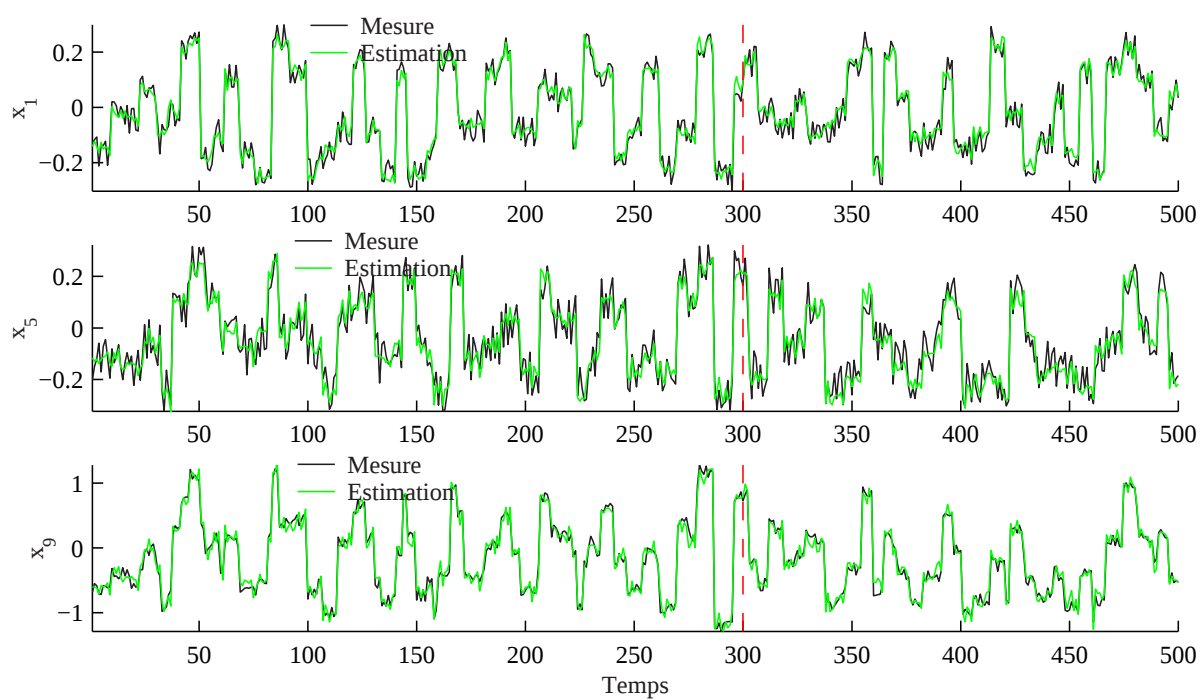


Figure 2.10 – Evolution des mesures et des estimations des trois variables  $x_1$ ,  $x_5$  et  $x_9$  de l'exemple 2



## 3

# Détection et localisation de défauts par Analyse en Composantes Principales

## Sommaire

|            |                                                      |           |
|------------|------------------------------------------------------|-----------|
| <b>3.1</b> | <b>Introduction</b>                                  | <b>33</b> |
| <b>3.2</b> | <b>Détection de défauts</b>                          | <b>35</b> |
| 3.2.1      | Génération de résidus par estimation d'état          | 35        |
| 3.2.2      | Génération de résidus par estimation paramétrique    | 56        |
| <b>3.3</b> | <b>Localisation de défauts</b>                       | <b>59</b> |
| 3.3.1      | Localisation par structuration des résidus           | 60        |
| 3.3.2      | Localisation utilisant un banc de modèles            | 69        |
| 3.3.3      | Localisation par calcul des contributions            | 85        |
| 3.3.4      | Localisation de défauts multiples                    | 92        |
| <b>3.4</b> | <b>Identification des caractéristiques du défaut</b> | <b>97</b> |
| <b>3.5</b> | <b>Conclusion</b>                                    | <b>98</b> |

## 3.1 Introduction

Récemment, les méthodes de détection et de localisation de défauts reposant sur l'analyse en composantes principales linéaires (ACP) ont reçu une attention particulière et ont été largement utilisées pour la surveillance des processus industriels [63, 87, 69, 98]. Le principe de cette approche est d'utiliser l'analyse en composantes principales pour modéliser le comportement du processus en fonctionnement normal et les défauts sont alors détectés en comparant le comportement observé et celui donné par le modèle ACP. La plupart des méthodes récentes utilisent l'erreur quadratique de prédiction  $SPE$  et la statistique de Hotelling  $T^2$  pour la détection de défauts sur les mesures [57, 55, 14, 97, 63, 113, 76, 27, 18]. Toutefois ces deux indices de détection jouent des rôles différents dans la stratégie de surveillance par ACP. La statistique  $T^2$  décrit le comportement des variables du processus qui sont corrélées avec les composantes principales tandis que la statistique  $SPE$  dépend de toutes les variables à surveiller. Cependant, le  $SPE$

est un test global qui cumule les erreurs de modélisation présentes sur chaque résidu [31] et la statistique  $T^2$  est calculée à partir des premières composantes principales qui ne représentent pas des résidus de plus que les conditions d'utilisation de cette statistique sont rarement vérifiées [4]. Pour améliorer les capacités de détection de la méthode d'analyse en composantes principales, un test basé sur les dernières composantes principales a été proposé [31].

Pour localiser la (ou les) variable en défaut, plusieurs méthodes de localisation de défauts utilisant l'analyse en composantes principales ont été proposées ces dix dernières années.

Inspirée des méthodes de localisation à base de redondances analytiques, la localisation de défauts utilisant la structuration des résidus à partir d'un modèle ACP a été récemment développée [25]. Une extension de cette approche, maximisant la sensibilité des résidus structurés aux défauts, a été proposée par Qin *et al.* [79]. Gertler *et al.* [26] utilisent une structuration particulière des résidus appelée ACP partielle. L'inconvénient majeur de cette approche est qu'il n'y a pas de méthode systématique permettant de choisir les ACP partielles. De plus, on est souvent confronté au problème de l'insensibilité des résidus à certains défauts, ce qui donne dans la plupart des cas des fausses localisations.

La méthode proposée par Dunia *et al.* [14] basée sur le principe de reconstruction est détaillée. Cette méthode est basée sur l'indice  $SPE$  pour la localisation. Elle suppose que chaque capteur peut être suspecté est reconstruit. Après la reconstruction de chaque variable un  $SPE$  est calculé. La comparaison du  $SPE$  avant et après reconstruction permet de définir la variable incriminée.

Une approche largement utilisée pour la localisation avec l'ACP consiste à calculer les contributions individuelles des variables à l'indice de détection ( $T^2$  ou  $SPE$ ) [55, 113, 68, 108]. La variable ayant la plus forte contribution à l'indice considéré est la variable incriminée.

Cependant, comme les deux statistiques  $T^2$  et  $SPE$  ne sont pas utilisés pour la détection, nous avons exploité les principes des deux méthodes de localisation utilisant le principe de reconstruction et le calcul des contributions pour les appliquer à l'indice de détection proposé.

Ainsi, concernant la méthode de localisation par reconstruction, le principe consiste à reconstruire chaque variable et de comparer l'indice proposé avant et après reconstruction. La reconstruction de la variable en défaut permet d'éliminer l'effet du défaut et l'indice de détection calculé après reconstruction de cette variable ne présente pas de dépassement de son seuil de détection. Ainsi, la variable incriminée peut être localisée par simple test sur les différents indices calculés après la reconstruction des différentes variables.

Pour la localisation en utilisant le calcul des contributions des variables à l'indice proposé, nous avons proposé deux définitions de ces contributions. La première définition exploite le fait que l'indice proposé est un  $SPE$  particulier calculer sur un sous-espace de l'espace résiduel et donc nous utilisons une définition similaire à celle donnée dans le cas du  $SPE$ . Puisque l'indice proposé peut être exprimé en fonction des dernières composantes, la seconde définition calcul les contributions des variables aux composantes intervenants dans le calcul de l'indice proposé ayant subi une variation significative due à la présence d'un défaut.

Ce chapitre a pour but de présenter le principe de détection et localisation de défauts basé sur l'analyse en composantes principales linéaires.

La deuxième section de ce chapitre présentera le principe de détection utilisant l'analyse en composantes principales. Ainsi, les deux approches de générations de résidus (estimation d'état et estimation paramétrique) seront exposées. Dans un premier temps nous présenterons les indices de détection utilisés dans le cas de l'approche par estimation d'état ainsi que l'indice de détection proposé. Ensuite une présentation de l'approche de détection par estimation paramétrique sera exposée.

Les différentes approches de localisation basée sur l'ACP seront présentées dans la troisième section. Dans le cadre de l'ACP, les deux approches les plus populaires, sont basées sur les prin-

cipes de reconstruction et des contributions aux indices de détection. En combinant la procédure de détection proposée avec le principe de reconstruction qui consiste à reconstruire une variable en utilisant le modèle ACP, une procédure de détection et de localisation utilisant les dernières composantes principales est proposée [33]. Une des méthodes de localisation qui est largement utilisée dans le cas de l'ACP est la méthode des contributions. Cette méthode est basée sur le calcul des contributions des différentes variables à l'indice de détection. L'application de cette approche à l'indice de détection proposé sera également détaillée dans cette section.

Les différentes approches exposées dans ce chapitre seront illustrées avec les deux exemples que nous avons présentés dans le chapitre 1.

## 3.2 Détection de défauts

Généralement, la phase de détection de défaut dans le cas du diagnostic à base de modèle analytique est liée à l'étape de génération de résidus qui a pour but de générer, à partir d'un modèle de bon fonctionnement du processus et des mesures disponibles, des signaux révélateurs de la présence de défauts, appelés résidus. A partir de l'analyse de ces résidus, l'étape de prise de décision doit alors indiquer si un défaut est présent ou non. Il existe deux approches pour la génération des résidus : l'approche par estimation d'état et l'approche par estimation des paramètres.

### 3.2.1 Génération de résidus par estimation d'état

Dans le cas de l'analyse en composantes principales on retrouve le même principe. Concernant l'approche par estimation d'état, les résidus sont obtenus par un test de cohérence entre l'estimation et les mesures.

La présence d'un défaut affectant l'une des variables provoque un changement dans les corrélations entre les variables indiquant une situation inhabituelle car les relations entre les variables ne sont plus vérifiées. Dans ce cas, la projection du vecteur de mesures dans le sous-espace des résidus va croître par rapport à sa valeur dans les conditions normales. Pour détecter un tel changement dans les corrélations entre les variables, l'ACP utilise généralement, la statistique *SPE* [5, 46] et la statistique  $T^2$  de Hotelling. Raich and Cinar [80] et Yue and Qin [115] ont proposé une statistique combinant le  $T^2$  et *SPE*. Cependant en pratique, la statistique *SPE* est généralement affectée par des erreurs de modélisation [31]. De plus, la statistique  $T^2$  ne représente par un indice de détection. Dans cette section, nous allons présenter les différents indices de détection utilisés avec l'ACP ainsi que les limitations de ces derniers et la solution que nous proposons pour résoudre le problème de détection en utilisant l'analyse en composantes principales.

Une fois le nombre de composantes  $\ell$  déterminé, la matrice  $X$  peut être approximée à partir des  $\ell$  premières composantes principales correspondant aux  $\ell$  plus grandes valeurs propres de la matrice  $\Sigma$  :

$$\hat{X} = \sum_{i=1}^{\ell} T_i p_i^T = \sum_{i=1}^{\ell} X p_i p_i^T \quad (3.1)$$

Notons que la matrice des vecteurs propres et la matrice des composantes principales peuvent chacune être décomposées en deux sous-matrices  $P = [\hat{P} \tilde{P}]$  et  $T = [\hat{T} \tilde{T}]$ . Les deux premières sous-matrices  $\hat{P}$  et  $\hat{T}$  représentent les matrices des  $\ell$  premiers vecteurs propres et les  $\ell$  premières composantes principales, respectivement. Les deux sous-matrices  $\tilde{P}$  et  $\tilde{T}$  représentent les matrices des  $(m - \ell)$  vecteurs propres et les  $(m - \ell)$  dernières composantes principales, respectivement.

Ainsi, la matrice des composantes principales  $T_{N \times m}$  est donnée par :

$$T = XP = X [\hat{P} \quad \tilde{P}] \quad (3.2)$$

Sachant que  $\Sigma = P\Lambda P^T$ , la matrice de covariance des données transformées  $\Sigma_t$  est donnée par :

$$\Sigma_t = P^T \Sigma P = P^T P \Lambda = \Lambda \quad (3.3)$$

A partir de l'équation (3.2), on peut écrire :

$$\hat{T} = X\hat{P} \quad (3.4)$$

et

$$\hat{X} = \hat{T}\hat{P}^T \quad (3.5)$$

$\hat{T}$  représente la projection de  $X$  sur les  $\ell$  premiers vecteurs propres de la matrice de covariance  $\Sigma$  et  $\tilde{T}$  représente la projection de  $X$  sur les  $(m - \ell)$  derniers vecteurs propres.

$$\tilde{T} = X\tilde{P} \quad (3.6)$$

et

$$\tilde{X} = \tilde{T}\tilde{P}^T \quad (3.7)$$

où  $\tilde{X}$  représente la matrice des résidus notée  $E$ .

On en déduit la décomposition suivante de la matrice  $X$  :

$$X = \hat{X} + \tilde{X} = \hat{X} + E \quad (3.8)$$

Les matrices  $\hat{X}$  et  $\tilde{X}$  représentent, respectivement, les variations modélisées et non modélisées de  $X$ . La matrice d'erreur est calculée par :

$$E = X\tilde{C} \quad (3.9)$$

et

$$\hat{X} = X\hat{C} \quad (3.10)$$

où  $\hat{C} = \hat{P}\hat{P}^T$  et  $\tilde{C} = (I - \hat{C})$ .

Pour un nouveau vecteur de mesure  $\mathbf{x}$  à un instant donné  $k$ , les équations précédentes deviennent :

$$\mathbf{x}(k) = \hat{\mathbf{x}}(k) + \tilde{\mathbf{x}}(k) = \hat{\mathbf{x}}(k) + \mathbf{e}(k) \quad (3.11)$$

Le vecteur des données transformées est donné par :

$$\mathbf{t}(k) = P^T \mathbf{x}(k) = [\hat{\mathbf{t}}(k) \quad \tilde{\mathbf{t}}(k)] \quad (3.12)$$

où  $\hat{\mathbf{t}}(k) = \hat{P}^T \mathbf{x}(k)$  et  $\tilde{\mathbf{t}}(k) = \tilde{P}^T \mathbf{x}(k)$

$$\hat{\mathbf{t}}(k) = [t_1, t_2, \dots, t_\ell]^T \quad (3.13)$$

et

$$\tilde{\mathbf{t}}(k) = [t_{\ell+1}, \dots, t_m]^T \quad (3.14)$$

$\hat{\mathbf{x}}(k) = \hat{C}\mathbf{x}(k)$  représente le vecteur des mesures estimées et  $\mathbf{e}(k) = (I - \hat{C})\mathbf{x}(k)$  représente le vecteur des erreurs d'estimation (fig 3.1). Cette relation est très intéressante pour la surveillance des processus, en effet, on peut calculer l'erreur quadratique  $SPE(k)$  (squared prediction error), connue aussi sous le nom de statistique  $Q$ , comme :

$$SPE(k) = \mathbf{e}(k)^T \mathbf{e}(k) \quad (3.15)$$

ou par :

$$SPE(k) = \|\tilde{\mathbf{x}}(k)\|^2 = \|\tilde{C}\mathbf{x}(k)\|^2 \quad (3.16)$$

où  $\tilde{C} = \tilde{P}\tilde{P}^T = (I - C)$ . Notons aussi que :

$$\mathbf{e}(k)^T \mathbf{e}(k) = \tilde{\mathbf{t}}^T \tilde{P}^T \tilde{P} \tilde{\mathbf{t}}(k) = \tilde{\mathbf{t}}(k)^T \tilde{\mathbf{t}}(k) \quad (3.17)$$

est une autre façon de calculer la quantité  $SPE(k)$ .

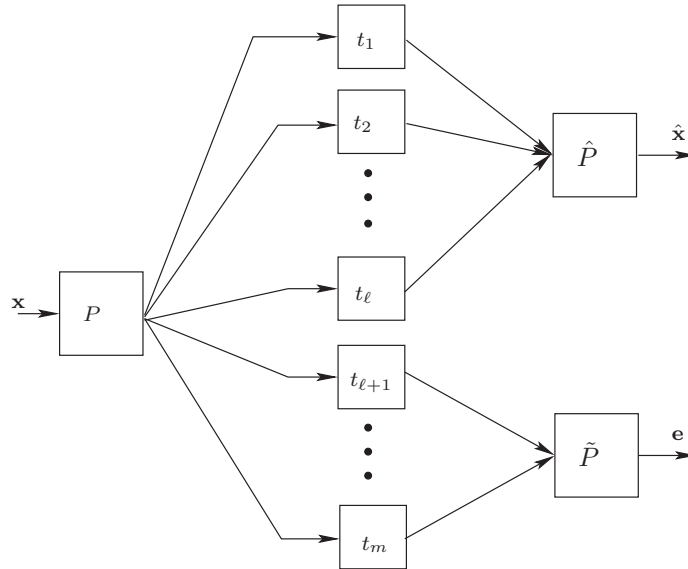


Figure 3.1 – Décomposition du vecteur de mesure  $\mathbf{x}$  en un vecteur de mesure estimé  $\hat{\mathbf{x}}$  et un vecteur de résidus  $\mathbf{e}$

Une interprétation géométrique de l'analyse en composantes principales peut être donnée (fig 3.2). Le modèle ACP partitionne l'espace  $\Re^m$  des mesures en un sous-espace des composantes principales,  $\mathcal{S}_{\hat{p}}$ , où les variations normales auront lieu et un sous-espace des résidus,  $\mathcal{S}_r$ , où les défauts doivent apparaître. Notons que  $\dim(\mathcal{S}_{\hat{p}}) = \ell$  et  $\dim(\mathcal{S}_r) = (m - \ell)$ . Donc, le vecteur des mesures peut être décomposé en deux parties :

$$\mathbf{x} = \hat{\mathbf{x}} + \tilde{\mathbf{x}} \quad (3.18)$$

où

$$\hat{\mathbf{x}} = C\mathbf{x} \in \mathcal{S}_{\hat{p}} \quad (3.19)$$

est la projection du vecteur dans le sous-espace des composantes principales, et

$$\tilde{\mathbf{x}} = (I - \hat{C})\mathbf{x} \in \mathcal{S}_r \quad (3.20)$$

est la projection dans le sous-espace des résidus.

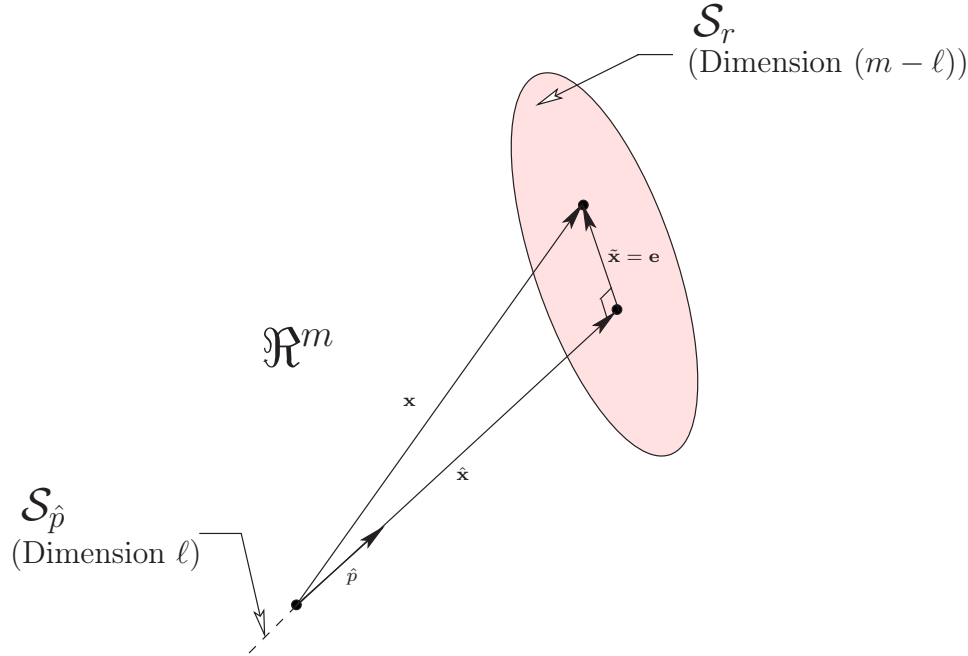


Figure 3.2 – Interprétation géométrique de l'analyse en composantes principales

### Dualité entre les relations de parité et l'ACP [25]

Soit le système linéaire statique avec les entrées  $\mathbf{u}(k) = [u_1(k) \dots u_a(k)]^T$  et les sorties  $\mathbf{y}(k) = [y_1(k) \dots y_b(k)]$  où  $a$  et  $b$  représentent le nombre de variables d'entrées et de sorties, respectivement.

Ces variables sont sujettes à des bruits de mesures  $\epsilon_y(k)$  et  $\epsilon_u(k)$ . Les valeurs observées de ces variables,  $\mathbf{u}(k)$  et  $\mathbf{y}(k)$ , sont reliées aux valeurs vraies  $\mathbf{u}^\circ$  et  $\mathbf{y}^\circ(k)$  par :

$$\begin{aligned} \mathbf{u}(k) &= \mathbf{u}^\circ(k) + \epsilon_u(k) \\ \mathbf{y}(k) &= \mathbf{y}^\circ(k) + \epsilon_y(k) \end{aligned} \quad (3.21)$$

Nous supposons que  $\mathbf{u}^\circ(k)$  et  $\mathbf{y}^\circ(k)$  sont reliées par l'équation :

$$\mathbf{y}^\circ(k) = A\mathbf{u}^\circ(k) \quad (3.22)$$

où la matrice  $A$  peut être considérée comme le modèle du système. Dans l'approche des relations de parité, la matrice  $A$  est supposée connue.

Pour illustrer le lien entre les relations de parité et l'ACP, nous allons combiner les vecteurs d'entrée et de sortie en un seul vecteur  $\mathbf{x}^*(k)$ , où l'étoile signifie que les mesures ne contiennent pas de défauts :

$$\mathbf{x}^*(k) = \begin{bmatrix} \mathbf{y}(k) \\ \mathbf{u}(k) \end{bmatrix} \quad (3.23)$$

et les erreurs de mesures associées :

$$\varepsilon^*(k) = \begin{bmatrix} \varepsilon_y(k) \\ \varepsilon_u(k) \end{bmatrix} \quad (3.24)$$

Ainsi, on peut écrire (3.21) et (3.22) comme :

$$\mathbf{x}^*(k) = \mathbf{x}^o(k) + \varepsilon^*(k) \quad (3.25)$$

$$B\mathbf{x}^*(k) = \varepsilon^*(k) \quad (3.26)$$

où  $B = [I - A]$  et  $\varepsilon(k)^*$  est le vecteur des résidus représentant les bruits de mesure.

Dans le cas de l'ACP, le vecteur des mesures, en fonctionnement normal, est donné par l'équation suivante :

$$\mathbf{x}^*(k) = \hat{P}\hat{\mathbf{t}}^*(k) + \tilde{P}\tilde{\mathbf{t}}^*(k) \quad (3.27)$$

Puisque  $\hat{P}$  et  $\tilde{P}$  sont orthogonaux, en prémultipliant cette équation par  $\tilde{P}^T$ , nous obtenons le vecteur des résidus suivant :

$$\tilde{P}^T \mathbf{x}^*(k) = \tilde{P}^T \tilde{P} \tilde{\mathbf{t}}^*(k) = \tilde{\mathbf{t}}^*(k) \quad (3.28)$$

qui représente le vecteur des dernières composantes principales.

Dans ce cas, en comparant les deux équations (3.26) et (3.28), nous obtenons :

$$B = \tilde{P}^T \quad (3.29)$$

Ainsi la matrice  $B$  permettant la génération de résidus dans le cadre des relations de parité est équivalente à la matrice  $\tilde{P}^T$  dans cas de l'ACP.

Avant d'aborder les différents indices de détection, nous allons montrer les expressions des résidus en présence d'un défaut  $d$  quelconque sur la  $j^{eme}$  variable du processus à surveiller. Notons par  $\mathbf{x}^*(k)$  le vecteur des mesures à l'instant  $k$  en absence de défaut et  $\xi_j$  la direction du défaut. Ainsi on peut écrire que :

$$\mathbf{x}(k) = \mathbf{x}^*(k) + \xi_j d(k) \quad (3.30)$$

A partir de cette équation le vecteur des résidus peut s'écrire sous la forme :

$$\begin{aligned} \mathbf{e}(k) &= \tilde{C}\mathbf{x}(k) \\ &= \begin{pmatrix} \tilde{c}_{11} & \tilde{c}_{12} & \dots & \tilde{c}_{1m} \\ \tilde{c}_{21} & . & . & . \\ . & . & . & . \\ . & . & . & . \\ \tilde{c}_{m1} & . & \dots & \tilde{c}_{mm} \end{pmatrix} (\mathbf{x}^*(k) + \xi_j d(k)) \\ &= \mathbf{e}^*(k) + \begin{pmatrix} \tilde{c}_{1j} \\ \tilde{c}_{2j} \\ . \\ . \\ \tilde{c}_{mj} \end{pmatrix} d(k) \end{aligned} \quad (3.31)$$

où  $\mathbf{e}^*$  représente le vecteur des résidus en absence de défaut qui sont statistiquement nuls. Ainsi, le défaut  $d$  se propage dans tout les résidus, et plus précisément un défaut  $d_j$  affectant la variable  $x_j$  se propage dans le résidu  $e_i$  avec une amplitude  $\tilde{c}_{ij}d_j$ .

De la même façon nous allons voir l'effet du défaut sur les composantes principales. Les premières composantes principales à l'instant  $k$  peuvent être données par l'équation suivante :

$$\begin{aligned}\hat{\mathbf{t}}(k) &= \hat{P}^T \mathbf{x}(k) \\ &= \begin{pmatrix} p_{11} & p_{12} & \dots & p_{1\ell} \\ p_{21} & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ p_{m1} & \cdot & \dots & p_{m\ell} \end{pmatrix}^T (\mathbf{x}^*(k) + \xi_j d(k)) \\ &= \hat{\mathbf{t}}^*(k) + \begin{pmatrix} p_{j1} \\ p_{j2} \\ \cdot \\ \cdot \\ p_{j\ell} \end{pmatrix} d(k)\end{aligned}\quad (3.32)$$

Ainsi, le défaut se propage dans toutes les composantes avec un amplitude  $p_{ji} d$  affectant la  $i^{eme}$  composante ( $i = 1, \dots, \ell$ ). Il faut noter que ce vecteur est statistiquement non nul.

De même pour les dernières composantes qui sont données par :

$$\begin{aligned}\tilde{\mathbf{t}}(k) &= \tilde{P}^T \mathbf{x}(k) \\ &= \begin{pmatrix} p_{1,\ell+1} & p_{1,\ell+2} & \dots & p_{1,m} \\ p_{2,\ell+1} & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot \\ p_{m,\ell+1} & \cdot & \dots & p_{m,m} \end{pmatrix}^T (\mathbf{x}^*(k) + \xi_j d(k)) \\ &= \tilde{\mathbf{t}}^*(k) + \begin{pmatrix} p_{j,\ell+1} \\ p_{j,\ell+2} \\ \cdot \\ \cdot \\ p_{jm} \end{pmatrix} d(k)\end{aligned}\quad (3.33)$$

La  $l^{eme}$  composante est affectée par un défaut d'amplitude  $p_{jl} d$  ( $l = \ell + 1, \dots, m$ ) et  $\tilde{\mathbf{t}}^*(k)$  est statistiquement nul.

### 3.2.1.1 Statistique SPE

Une statistique typique pour détecter ces conditions anormales est la statistique *SPE*, appelée aussi *Q* (Squared Prediction Error) qui est donnée par l'équation :

$$SPE(k) = \mathbf{e}(k)^T \mathbf{e}(k) \quad (3.34)$$

Le processus est considéré en fonctionnement anormal (présence d'un défaut) à l'instant  $k$  si :

$$SPE(k) > \delta_\alpha^2 \quad (3.35)$$

où  $\delta^2$  est le seuil de confiance du  $SPE(k)$ . Soit  $\theta_i = \sum_{j=\ell+1}^m \lambda_j^i$  pour  $i = 1, 2, 3$  et  $\lambda_j$  est la  $j^{eme}$  valeur propre de la matrice  $\Sigma$ .

Le seuil est approximé par [46] :

$$\delta_\alpha^2 = \theta_1 \left[ \frac{c_\alpha \sqrt{2\theta_2 h_0^2}}{\theta_1} + 1 + \frac{\theta_2 h_0 (h_0 - 1)}{\theta_1^2} \right]^{\frac{1}{h_0}} \quad (3.36)$$

$$\text{où } h_0 = 1 - \frac{2\theta_1\theta_3}{3\theta_2^2} \text{ et } c_\alpha = \frac{\theta_1 \left[ \left( \frac{\|\mathbf{e}\|^2}{\theta_1} \right)^{h_0} - 1 - \frac{\theta_2 h_0 (h_0 - 1)}{\theta_1^2} \right]}{\sqrt{2\theta_2 h_0^2}}$$

$c_\alpha$  est la limite au seuil de confiance  $(1 - \alpha)$  dans le cas d'une distribution normale. Il faut noter que ce résultat est donné sous les conditions suivantes :

1. le vecteur des échantillons  $\mathbf{x}$  suit une distribution normale multivariable,
2. une approximation est effectuée pour calculer le seuil  $\delta_\alpha^2$  du  $SPE$ .

Box [5] a antérieurement montré que :

$$\delta_\alpha^2 = g\chi_{h,\alpha}^2 \quad (3.37)$$

où  $g = \theta_2/\theta_1$  et  $h = \theta_1^2/\theta_2$ . Plus récemment, Nomikos et MacGregor [71] ont démontré que les deux approximations données par les équations (3.36) et (3.37) sont équivalentes.

### Conditions de détectabilité avec le $SPE$

Le vecteur des mesures en présence d'un défaut  $d$  affectant la  $j^{eme}$  variable est donnée par :

$$\mathbf{x}(k) = \mathbf{x}^*(k) + \xi_j d(k) \quad (3.38)$$

Puisque  $\xi_j \in \mathbb{R}^m$ , il peut être projeté sur  $\mathcal{S}_p$  (sous-espace des composantes principales) et  $\mathcal{S}_r$  (sous-espace des résidus) comme :

$$\xi_j = \hat{\xi}_j + \tilde{\xi}_j \quad (3.39)$$

où  $\hat{\xi}_j = \hat{C}\xi_j \in \mathcal{S}_p$  et  $\tilde{\xi}_j = \tilde{C}\xi_j \in \mathcal{S}_r$ .

$$\mathbf{e}(k) = \mathbf{e}^*(k) + \tilde{\xi}_j d(k) \quad (3.40)$$

Ainsi, si on a un défaut pour lequel  $\tilde{\xi}_j d(k) \neq 0$ , le  $SPE(k)$  croît, ce qui indique la présence d'un défaut. Toutefois, pour garantir que le défaut soit détectable, une condition est rajoutée sur l'amplitude du défaut :

$$\tilde{d}_j(k) \leq \|\mathbf{e}(k)\| + \|\mathbf{e}^*(k)\| \quad (3.41)$$

où  $\tilde{d}_j(k) = d(k) \|\tilde{\xi}_j\|$ .

Puisque  $\|\mathbf{e}^*(k)\|^2$  représente le  $SPE(k)$  dans les conditions normales,  $\|\mathbf{e}^*(k)\| \leq \delta$  définit la région de confiance. Donc on peut écrire :

$$\|\mathbf{e}(k)\| \geq |\tilde{d}(k)| - \delta \quad (3.42)$$

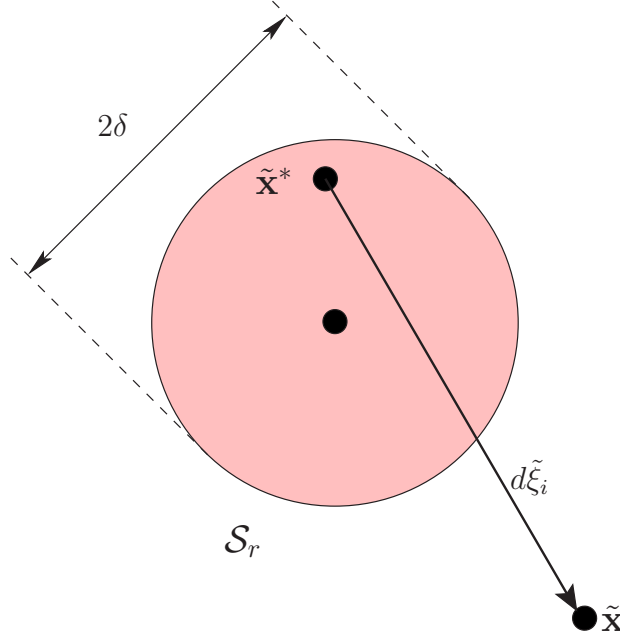


Figure 3.3 – Illustration graphique de la condition de détectabilité avec le *SPE*

Pour garantir la détectabilité du défaut avec la condition :  $SPE(k) = \|\mathbf{e}(k)\|^2 > \delta^2$ , il faut que (fig 3.3) :

$$|\tilde{d}(k)| > 2\delta \quad (3.43)$$

### 3.2.1.2 Statistique $T^2$ de Hotelling

La statistique  $T^2$  peut être appliquée, dans le cas de l'analyse en composantes principales, sur les premières composantes principales. Ainsi, on obtient :

$$T^2(k) = \hat{\mathbf{t}}^T(k) \Lambda_\ell^{-1} \hat{\mathbf{t}}(k) \quad (3.44)$$

Cette quantité suit une distribution du Chi-2 ( $\chi^2$ ) avec  $\ell$  degrés de liberté où  $\Lambda_\ell = \text{diag}\{\lambda_1, \lambda_2, \dots, \lambda_\ell\}$  est une matrice diagonale contenant les  $\ell$  plus grandes valeurs propres de la matrice de corrélation.

Le seuil approprié pour la statistique  $T^2$  pour un seuil de confiance  $\alpha$  peut être déterminé. La statistique  $\frac{N(N-\ell)}{\ell(N+1)(N-1)} T^2$  suit une distribution de Fisher  $\mathbf{F}_{\ell, (N-\ell)}$ . La limite supérieure pour un seuil de confiance  $\alpha$  est obtenue à partir de  $\mathbf{F}_{\ell, (N-\ell), \alpha}$ , soit :

$$\chi_{\ell, \alpha}^2 = \frac{\ell(N+1)(N-1)}{N(N-\ell)} \mathbf{F}_{\ell, (N-\ell), \alpha} \quad (3.45)$$

Le processus est supposé en défaut (déviation anormale), à l'instant  $k$ , si :

$$T^2(k) > \chi_{\ell, \alpha}^2 \quad (3.46)$$

Puisque la statistique  $T^2$  n'est pas affectée par le bruit, qui est représenté par les dernières valeurs propres, théoriquement elle est capable de représenter le comportement normal du processus. La statistique  $T^2$  peut être interprétée comme la mesure des variations normales du processus, et la violation du seuil de détection de cette statistique indique que ces variations sont en dehors des limites de contrôle et correspondent à un fonctionnement anormal.

### Condition de détectabilité avec le $T^2$

De la même façon que dans le cas du  $SPE$ , la condition suffisante de détectabilité avec le  $T^2$  peut être obtenue à partir de l'équation (3.44) :

$$\begin{aligned} T^2(k) &= \left\| \Lambda_\ell^{-1/2} \hat{P}^T \mathbf{x}(k) \right\|^2 = \left\| \Lambda_\ell^{-1/2} \hat{P}^T \hat{P} \hat{P}^T \mathbf{x}(k) \right\|^2 \\ &= \left\| \Lambda_\ell^{-1/2} \hat{P}^T \left( \hat{P} \hat{P}^T \mathbf{x}^*(k) + \hat{P} \hat{P}^T \xi_i d(k) \right) \right\|^2 \end{aligned} \quad (3.47)$$

$$T^2(k) = \left\| \Lambda_\ell^{-1/2} \hat{P}^T \left( \hat{\mathbf{x}}^*(k) + \hat{\xi}_i d(k) \right) \right\|^2 \geq \left\| \Lambda_\ell^{-1/2} \hat{P}^T \hat{\xi}_i d(k) \right\|^2 - \left\| \Lambda_\ell^{-1/2} \hat{P}^T \hat{\mathbf{x}}^*(k) \right\|^2 \quad (3.48)$$

Puisque  $\left\| \Lambda_\ell^{-1/2} \hat{P}^T \hat{\mathbf{x}}^*(k) \right\|^2 \leq \chi_{\ell, \alpha}^2$ , pour garantir que  $T^2(k) > \chi_{\ell, \alpha}^2$ , il faut avoir :

$$\left\| \Lambda_\ell^{-1/2} \hat{P}^T \hat{\xi}_i d(k) \right\| \geq 2\chi_{\ell, \alpha} \quad (3.49)$$

#### 3.2.1.3 Statistique SWE

Un autre indice de détection  $SWE$  (Squared Weighted Error), plus sensible au défaut [109], peut être défini comme le  $SPE$  pondéré par l'inverse de la variance des dernières composantes, et n'est en fait que la statistique  $T^2$  appliquée aux dernières composantes :

$$SWE(k) = \mathbf{x}^T(k) \tilde{P} \Lambda_{(m-\ell)}^{-1} \tilde{P}^T \mathbf{x}(k) \quad (3.50)$$

où  $\Lambda_{(m-\ell)} = \text{diag} \{ \lambda_{\ell+1}, \dots, \lambda_m \}$  est une matrice diagonale contenant les  $(m-\ell)$  dernières valeurs propres de la matrice de corrélation  $\Sigma$ .

Cet indice suit une distribution du chi-2 avec  $(m-\ell)$  degrés de liberté :

$$SWE(k) = \tilde{\mathbf{t}}^T(k) \Lambda_{(m-\ell)}^{-1} \tilde{\mathbf{t}}(k) \sim \chi_{(m-\ell), \alpha}^2 \quad (3.51)$$

Le processus est supposé en défaut si :

$$SWE(k) > \chi_{(m-\ell), \alpha}^2 \quad (3.52)$$

Il est important de noter que cet indice n'impose aucune supposition sur la distribution des variables du processus  $\mathbf{x}$ , mais il suppose que les résidus  $\tilde{\mathbf{t}}(k) = \tilde{P}^T \mathbf{x}(k)$  sont des bruits blancs gaussiens centrés.

### 3.2.1.4 Statistique combinée

La statistique combinée utilise à la fois les deux statistiques  $T^2$  et  $SPE$  pour obtenir un nouveau indice de détection de défauts. Pour expliquer l'idée d'un tel indice, nous allons considérer un exemple à deux composantes, ainsi la statistique  $T^2$  sera donnée par :

$$T^2(k) = \mathbf{x}^T(k) \hat{P} \Lambda_2^{-1} \hat{P}^T \mathbf{x}(k) \leq \chi_{2,\alpha}^2 \quad (3.53)$$

avec

$$\Lambda_2 = \begin{pmatrix} \lambda_1 & 0 \\ 0 & \lambda_2 \end{pmatrix} \quad (3.54)$$

on peut écrire :

$$T^2 = \frac{t_1^2(k)}{\lambda_1} + \frac{t_2^2(k)}{\lambda_2} \leq \chi_{2,\alpha}^2 \quad (3.55)$$

Cette équation représente l'équation d'une ellipse. La figure (fig 3.4) présente une illustration graphique de la détection de défaut en utilisant les deux statistiques  $SPE$  et  $T^2$ .

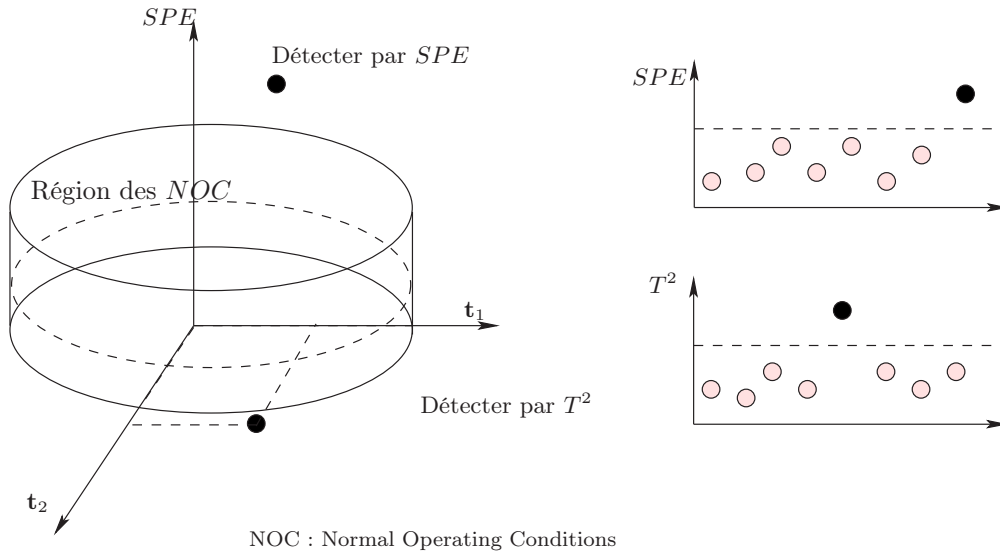


Figure 3.4 – Illustration graphique de la détection de défaut par les deux statistiques  $SPE$  et  $T^2$

Les deux statistiques permettent de déterminer une région de fonctionnement normal limitée par l'ellipse que définit le seuil  $\chi_\alpha^2$  de la statistique  $T^2$  et le seuil  $\delta_\alpha^2$  du  $SPE$ . Théoriquement, cette zone représente, pour chaque échantillon, la zone de fonctionnement normal et toute mesure qui se trouve à l'extérieur de cette zone est jugée en défaut.

Yue and Qin [115] proposent la formulation suivante de ce nouvel indice :

$$\zeta(k) = \frac{SPE(k)}{\delta_\alpha^2} + \frac{T^2(k)}{\chi_{\ell,\alpha}^2} = \mathbf{x}^T(k) \mathcal{M} \mathbf{x}(k) \quad (3.56)$$

où  $\mathcal{M}$  est donné par :

$$\mathcal{M} = \frac{(I - \hat{P}\hat{P}^T)}{\delta_\alpha^2} + \frac{\hat{P}\Lambda_\ell^{-1}\hat{P}^T}{\chi_{\ell,\alpha}^2} = \frac{\tilde{P}\tilde{P}^T}{\delta_\alpha^2} + \frac{\hat{P}\Lambda_\ell^{-1}\hat{P}^T}{\chi_{\ell,\alpha}^2} \quad (3.57)$$

Notons que  $\mathcal{M}$  est une matrice symétrique et définie positive.

$$\zeta(k) = \mathbf{x}^T(k) \mathcal{M} \mathbf{x}(k) \sim g \chi_h^2 \quad (3.58)$$

où le coefficient  $g$  est donné par :

$$g = \frac{\sum_{j=1}^m \lambda_j^2}{\sum_{j=1}^m \lambda_j} = \frac{\text{trace}(\Sigma \mathcal{M})^2}{\text{trace}(\Sigma \mathcal{M})} \quad (3.59)$$

$\lambda_j$  est la  $j^{\text{eme}}$  valeur propre de la matrice  $\Sigma \mathcal{M}$  et le nombre de degrés de liberté pour le  $\chi^2$  est donné par :

$$h = \frac{\left( \sum_{j=1}^m \lambda_j \right)^2}{\sum_{j=1}^m \lambda_j^2} = \frac{(\text{trace}(\Sigma \mathcal{M}))^2}{\text{trace}(\Sigma \mathcal{M})^2} \quad (3.60)$$

Après calcul de  $g$  et  $h$ , le défaut peut être détecté si :

$$\zeta(k) \geq g \chi_{h,\alpha}^2 \quad (3.61)$$

Il faut noter que cette statistique a été proposée pour la première fois par Raich et Cinar [80] sous une forme un peu différente que l'on notera  $v$  :

$$v(k) = \varsigma \frac{SPE(k)}{\delta_\alpha^2} + (1 - \varsigma) \frac{T^2(k)}{\chi_{\ell,\alpha}^2} \quad (3.62)$$

où  $\varsigma \in [0, 1]$  est une constante. Cette statistique indique un fonctionnement normal si elle est inférieure à 1, sinon la présence d'un défaut est suspectée.

### 3.2.1.5 Filtrage EWMA pour la détection

Pour améliorer la qualité de la détection et réduire le taux de fausses alarmes, le filtre EWMA (exponentially weighted moving average) peut être appliqué aux résidus. L'expression générale de ce filtre appliqué aux résidus est donnée par :

$$\bar{\mathbf{e}}(k) = (I - \beta) \bar{\mathbf{e}}(k-1) + \beta \mathbf{e}(k) \quad (3.63)$$

où  $\beta$  est une matrice diagonale dont les éléments sont les facteurs d'oubli pour les résidus,  $I$  est une matrice identité et  $\bar{\mathbf{e}}(0) = 0$ .

$$\overline{SPE}(k) = \|\bar{\mathbf{e}}(k)\|^2 \quad (3.64)$$

où  $\bar{\mathbf{e}}(k)$  et  $\overline{SPE}(k)$  sont les vecteurs des résidus et  $SPE$  filtrés, respectivement.  $\beta$  peut être ajusté en fonction du type de défauts à détecter :  $\beta$  proche de la matrice identité favorise la détection des changements lents, tandis que  $\beta$  proche de zéro est plus sensible aux changements brusques.

Le  $\overline{SPE}$  permet de réduire les fausses alarmes mais il introduit un certain retard à la détection.

La statistique  $Q$  originale développée par Jackson [46] n'est pas applicable au  $\overline{SPE}$ . Un test statistique pour le  $\overline{SPE}$  a été développé par Qin [76]. Pour simplifier, nous supposons dans la suite que la matrice des facteurs d'oubli est donnée par :

$$\beta = \gamma I \quad (3.65)$$

où  $\gamma$  est un facteur d'oubli. L'équation (3.63) équivaut à filtrer les données, puis à calculer les résidus. D'où :

$$\bar{\mathbf{x}}(k) = (1 - \gamma) \bar{\mathbf{x}}(k - 1) + \gamma \mathbf{x}(k) \quad (3.66)$$

$$\bar{\mathbf{e}}(k) = (I - C) \bar{\mathbf{x}}(k) \quad (3.67)$$

où  $\bar{\mathbf{x}}(k)$  représente ici le vecteur des données filtrées. En supposant que  $\mathbf{x}(k)$  est un processus aléatoire indépendant, la relation de covariance suivante peut être obtenue à partir de l'équation (3.66) :

$$\begin{aligned} \bar{\Sigma} &= \mathcal{E} \left\{ \bar{\mathbf{x}}(k) \bar{\mathbf{x}}^T(k) \right\} = \mathcal{E} \left\{ ((1 - \gamma) \bar{\mathbf{x}}(k - 1) + \gamma \mathbf{x}(k)) ((1 - \gamma) \bar{\mathbf{x}}^T(k - 1) + \gamma \mathbf{x}^T(k)) \right\} \\ &= (1 - \gamma)^2 \mathcal{E} \left\{ \bar{\mathbf{x}}(k - 1) \bar{\mathbf{x}}^T(k - 1) \right\} + \gamma^2 \mathcal{E} \left\{ \mathbf{x}(k) \mathbf{x}^T(k) \right\} \\ &\quad + 2(1 - \gamma) \mathcal{E} \left\{ \mathbf{x}(k) \bar{\mathbf{x}}^T(k - 1) \right\} \\ &= (1 - 2\gamma + \gamma^2) \bar{\Sigma} + \gamma^2 \Sigma \end{aligned} \quad (3.68)$$

$$\bar{\Sigma} = \frac{\gamma}{2 - \gamma} \Sigma \quad (3.69)$$

où  $\Sigma$  est la matrice de covariance de  $\mathbf{x}(k)$  et  $\bar{\Sigma}$  la matrice de covariance de  $\bar{\mathbf{x}}(k)$ . Comme conséquence, les valeurs propres de  $\bar{\Sigma}$  et  $\Sigma$  sont reliées par :

$$\bar{\lambda}_i = \frac{\gamma}{2 - \gamma} \lambda_i \quad \text{pour } i = 1, \dots, m \quad (3.70)$$

Donc, un raisonnement similaire à celui de Jackson [46], permet d'écrire :

$$\left[ \frac{\|\bar{\mathbf{e}}(k)\|^2}{\bar{\theta}_1} \right]^{\bar{h}} \sim \mathcal{N} \left[ 1 + \bar{\theta}_2 \bar{h}_0 \frac{(\bar{h}_0 - 1)}{\bar{\theta}_1^2}, \frac{2\bar{\theta}_2 \bar{h}_0^2}{\bar{\theta}_1^2} \right] \quad (3.71)$$

$$\bar{\theta}_i = \sum_{j=\ell+1}^m \bar{\lambda}_j^i = \left( \frac{\gamma}{2-\gamma} \right)^i \sum_{j=\ell+1}^m \lambda_j^i = \left( \frac{\gamma}{2-\gamma} \right)^i \theta_i \quad \text{pour } i = 1, 2, 3.$$

puisque,  $\bar{h}_0 = 1 - \frac{2\bar{\theta}_1 \bar{\theta}_3}{3\bar{\theta}_3^2} = 1 - \frac{2\theta_1 \theta_3}{3\theta_3^2} = h_0$ , la distribution dans l'équation (3.71) peut être simplifiée comme :

$$\left[ \frac{\|\bar{\mathbf{e}}(k)\|^2}{\bar{\theta}_1} \right]^{\bar{h}} \sim \mathcal{N} \left[ 1 + \frac{\theta_2 h_0 (h_0 - 1)}{\theta_1^2}, \frac{2\theta_2 h_0^2}{\theta_1^2} \right] \quad (3.72)$$

La statistique *SPE* filtrée est :

$$\overline{SPE}(k) > \bar{\delta}_\alpha^2 \quad (3.73)$$

$$\bar{\delta}_\alpha^2 = \bar{\theta}_1 \left[ \frac{c_\alpha \sqrt{2\theta_2 h_0^2}}{\theta_1} + 1 + \frac{\theta_2 h_0 (h_0 - 1)}{\theta_1^2} \right]^{\frac{1}{h_0}} \quad (3.74)$$

En utilisant la relation de l'équation (3.70) et en comparant les équations (3.36) et (3.74), nous obtenons la relation suivante :

$$\bar{\delta}_\alpha^2 = \frac{\gamma}{2 - \gamma} \delta_\alpha^2 \quad (3.75)$$

Cette équation relie la statistique  $SPE$  filtrée à celle du  $SPE$  non filtrée par une constante. Le  $\overline{SPE}$  définit une région de confiance plus étroite que celle du  $SPE$  à cause du filtrage.

Il faut noter que le filtrage peut être appliqué aux autres indices de détection. Nous l'avons présenté uniquement dans le cas du  $SPE$  pour des raisons que nous verrons ultérieurement.

Pour montrer l'intérêt du filtrage, nous allons considérer l'exemple 1.

Sur la figure (fig 3.5), on présente l'évolution du  $SPE$  et  $\overline{SPE}$  obtenue à partir de l'exemple 1 avec un défaut affectant la variable  $x_3$  à partir de l'instant 300 avec une amplitude qui s'élève à environ 20% de la plage de variation de cette variable. Le filtrage permet d'améliorer la détection, mais il introduit un certain retard à la détection. La figure (fig 3.6) présente l'évolution de  $T^2$  avec le même défaut sur la variable  $x_3$ . A partir de cette figure, on peut constater que le défaut n'est pas détectable sur  $T^2$ .

Le tableau (tab 3.1) présente les résultats de détection avec l'indice  $SPE$  filtré (fig 3.7) pour différents défauts simulés sur différentes variables de l'exemple 1.  $d_i$  représente le défaut sur la  $i^{eme}$  variable et  $k_d$  est l'instant d'apparition du défaut.

| Nature du défaut               | Biais        | Dérive                 | Défaillance complète |
|--------------------------------|--------------|------------------------|----------------------|
| Variable en défaut             | $x_1$        | $x_4$                  | $x_7$                |
| Expression du défaut           | $d_1 = 0.13$ | $d_4 = 0.005(k - k_d)$ | $x_7 = 0.3$          |
| Indice de détection            | $SPE$ Filtré | $SPE$ Filtré           | $SPE$ Filtré         |
| Instant d'apparition du défaut | 350          | 300                    | 320                  |
| Instant de détection           | 352          | 314                    | 323                  |

Table 3.1 – Simulation de défauts et résultats de détection sur l'exemple 1 <sup>a</sup>

<sup>a</sup>Les défauts simulés représentent, respectivement : un biais  $d_1$  d'une amplitude de 20% de la plage de variation de la variable  $x_1$ , une dérive  $d_4$  variant entre 0.8% et 150% et qui a été détectée une fois atteint une amplitude de 11% de la plage de variation de  $x_4$ , une défaillance complète  $d_7$  d'amplitude 23% de la plage de variation de  $x_7$ .

Il faut noter qu'aucun de ces défauts n'a été détecté par la statistique  $T^2$ . En réalité, la statistique  $T^2$  telle qu'elle a été définie dans le cas de l'ACP ne peut être utilisée comme un indice pour la détection de défaut. Nous constatons en fait que les premières composantes principales, qui interviennent dans le calcul de la statistique  $T^2$ , ne sont pas des résidus. Les résidus sont représentés par les dernières composantes principales comme le montre la figure (fig 2.3) représentant l'évolution des composantes principales de l'exemple 1.

Concernant l'indice  $SWE$  et d'après l'équation (3.51), cet indice n'est défini que si les dernières valeurs propres sont non nulles. Ainsi, si les valeurs propres sont nulles voire de faibles valeurs (ce qui est souvent le cas), cet indice n'est pas défini.

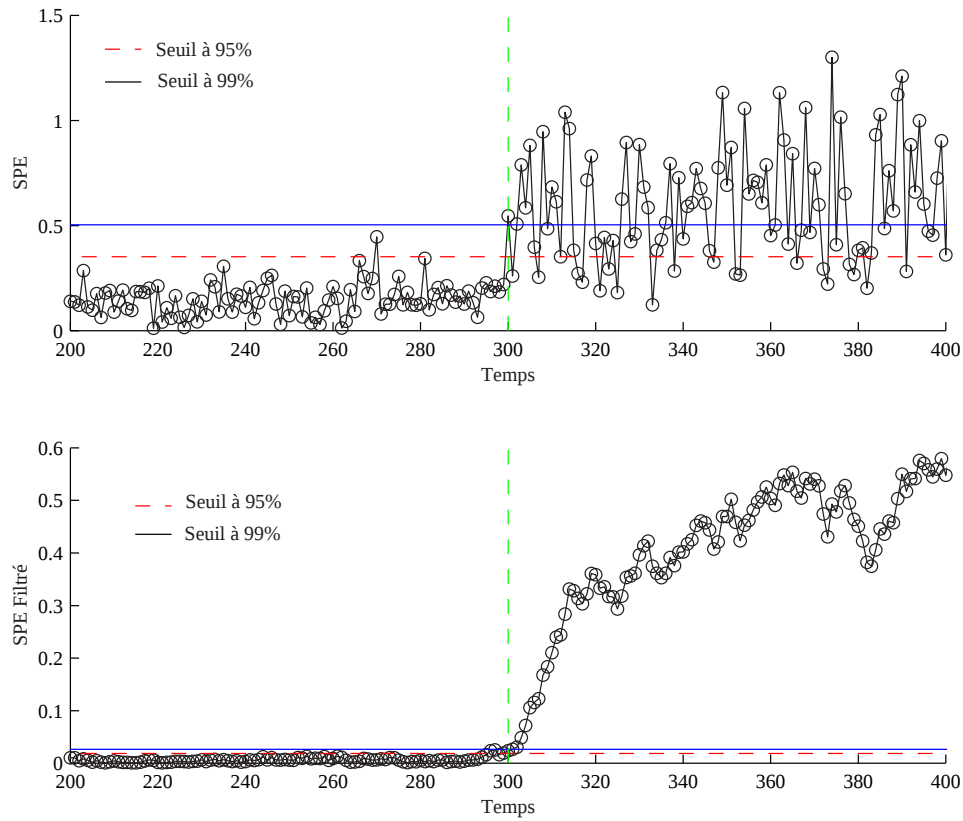


Figure 3.5 – Evolution du  $SPE$  et du  $SPE$  filtré avec un défaut affectant  $x_3$  à partir de l'instant 300 (exemple 1)

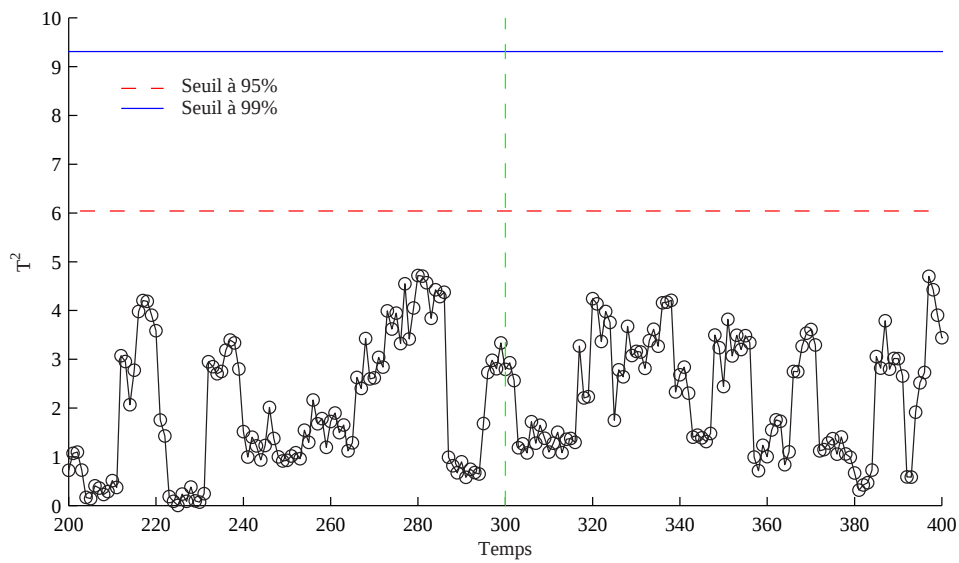


Figure 3.6 – Evolution de la statistique  $T^2$  avec un défaut affectant  $x_3$  à partir de l'instant 300.

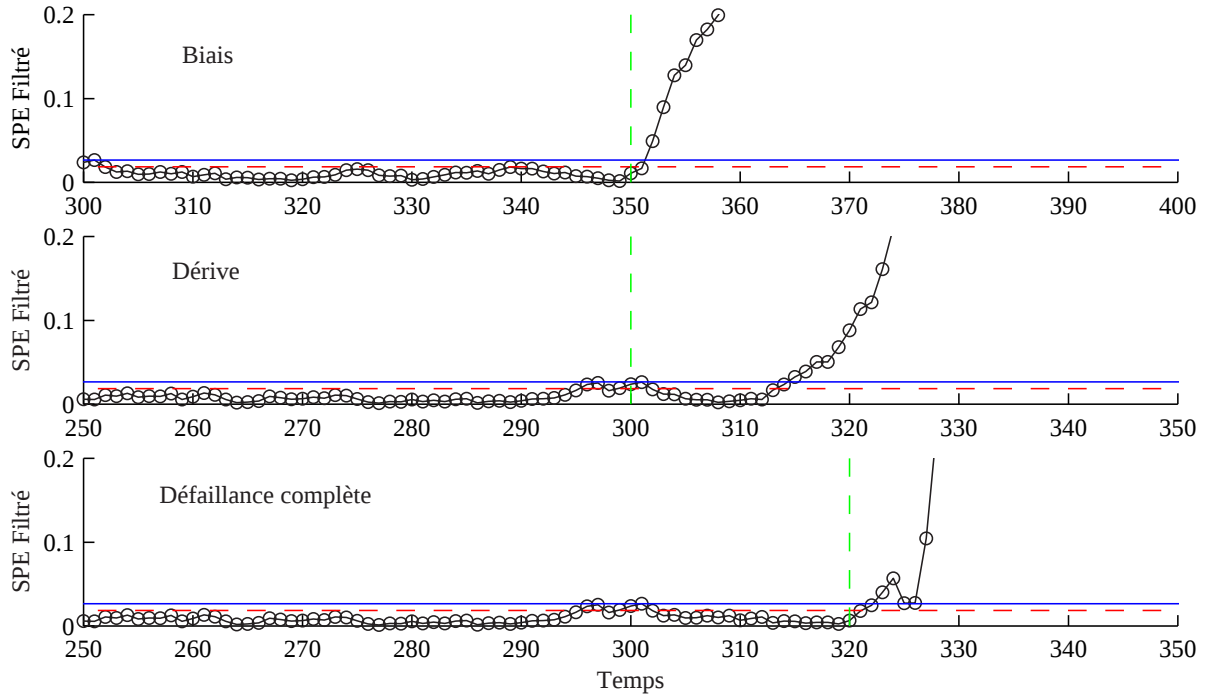


Figure 3.7 – Evolution des  $SPE$  filtrés correspondant aux différents défauts simulés et représentés sur le tableau (tab 3.1).

De plus, considérons l'exemple 2, la figure (fig 3.8) présente l'évolution de  $SPE$  et  $SPE$  filtré, calculés à partir d'un modèle ACP à trois composantes, en absence de défaut.

On constate qu'en absence de défaut la détection à partir du  $SPE$  filtré entraîne de nombreuses fausses alarmes. Si on augmente le seuil de façon à ne plus être sensible à ces erreurs, on risque de ne plus détecter certains défauts. Pour illustrer ce cas, nous allons simuler un défaut affectant la variable  $x_3$  entre les instants 300 et 500 avec une amplitude d'environ 20% de la plage de variation de cette variable.

A partir des figures (fig 3.9) et (fig 3.10), on constate que le défaut n'est pas détecté sur le  $SPE$  et que le fait de filtrer n'améliore pas la situation car le  $SPE$  filtré présente de nombreuses fausses alarmes.

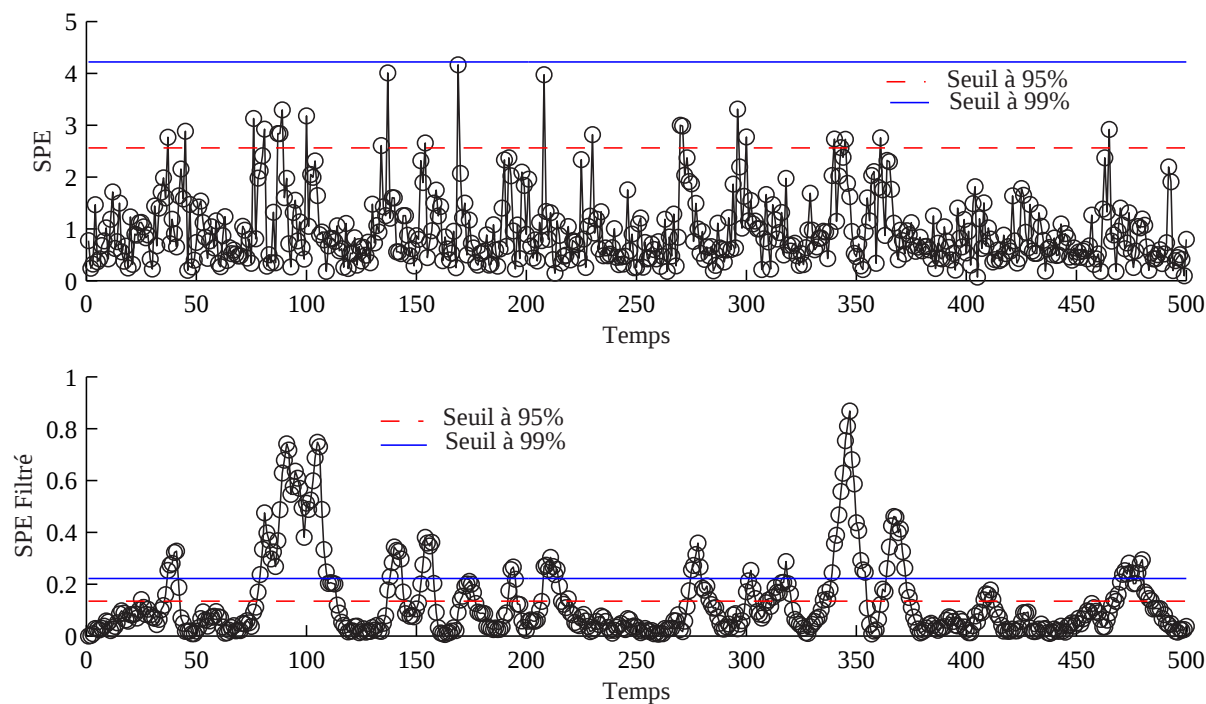
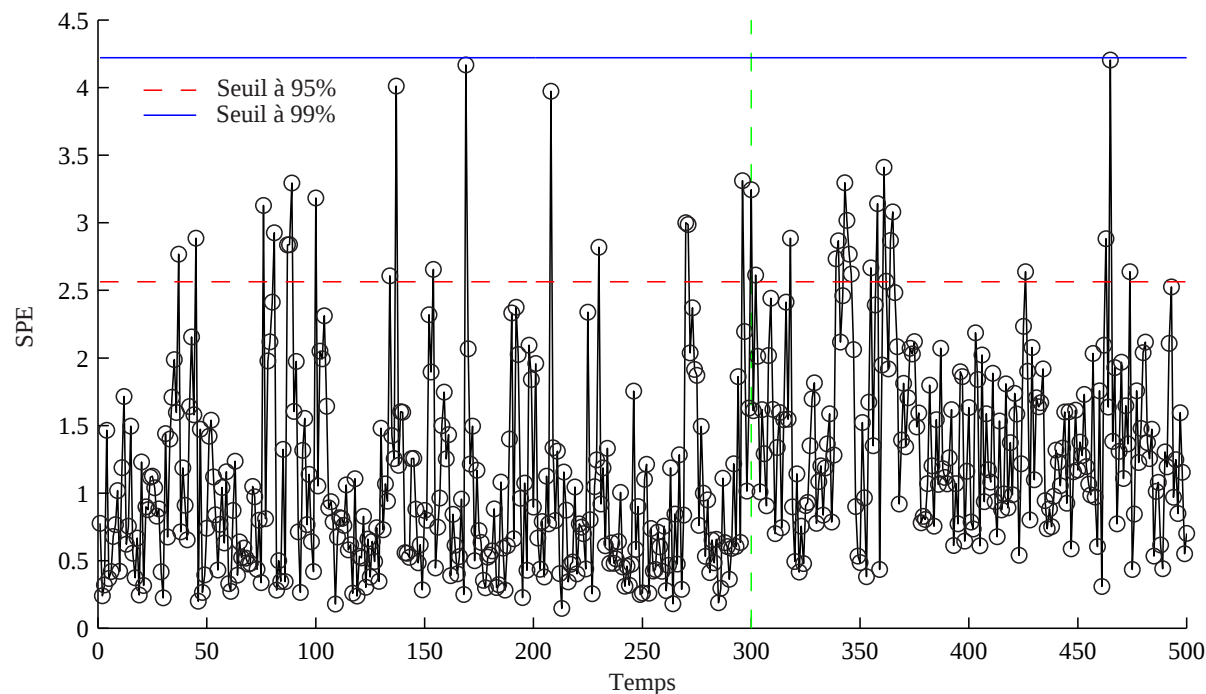
En comparant les figures (fig 3.5), (fig 3.8) et (fig 3.9), on constate qu'en absence d'erreurs de modélisation de  $SPE$  filtré permet d'améliorer la qualité de la détection par rapport au  $SPE$  (fig 3.5) et qu'en présence d'erreurs de modélisation (cas de l'exemple 2) le  $SPE$  filtré présente de nombreuses fausses alarmes qui sont dues à ces erreurs.

En conclusion, nous dirons que le  $SPE$  n'est pas un test efficace pour la détection du fait que c'est un test global qui cumule toutes les erreurs des résidus (3.34) et que le  $T^2$  ne représente pas un indice de détection d'autant plus que les conditions de son utilisation son rarement vérifiées, ce qui implique que l'indice combiné ne peut être exploité.

D'où la nécessité de développer un nouveau indice de détection.

### 3.2.1.6 Statistique $D$ proposée

Pour résoudre le problème de la détection de défaut, on suggère d'utiliser un test basé sur la somme des carrés des dernières composantes principales, que l'on note  $D_i$  ( $i = 1, 2, \dots, (m - \ell)$ )

Figure 3.8 – Evolution du  $SPE$  et  $SPE$  filtré en absence de défaut (cas de l'exemple 2)Figure 3.9 – Evolution du  $SPE$  avec un défaut affectant la variable  $x_3$  à partir de l'instant 300

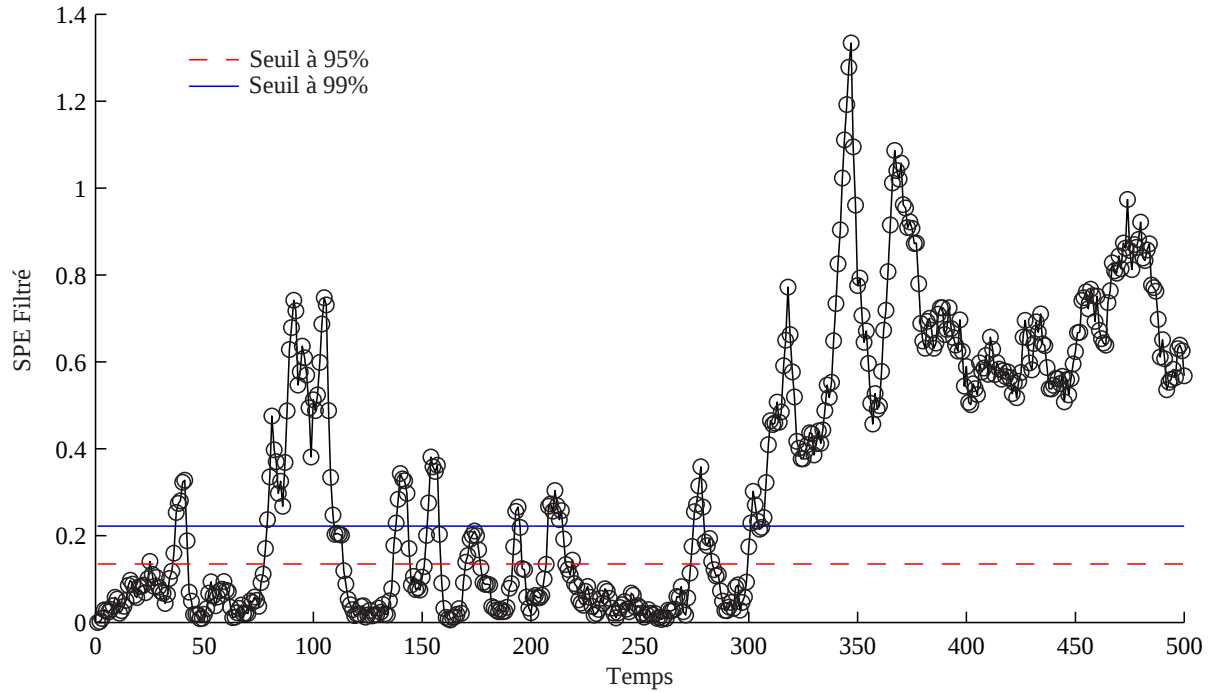


Figure 3.10 – Evolution du  $\overline{SPE}$  avec un défaut affectant la variable  $x_3$  à partir de l'instant 300

et qui est calculé de la façon suivante :

$$D_i(k) = \sum_{j=m-i+1}^m t_j^2(k), \quad i = 1, 2, \dots, (m - \ell) \quad (3.76)$$

Ainsi, le  $D_i(k)$  est un indice de détection de défaut calculé à partir des dernières composantes principales  $\tilde{\mathbf{t}}(k)$  à la différence de l'indice  $T^2$  qui est calculé à partir des premières composantes.

En analysant l'expression de  $D_i$  et en tenant compte du fait que  $\sum_{j=\ell+1}^m t_j^2(k) = SPE(k)$  est calculé avec un modèle ACP à  $\ell$  composantes, l'indice  $D_i$  correspond à un  $SPE$  calculé avec un modèle ACP à  $(m - i)$  composantes principales.

De ce fait, les seuils de détection  $\tau_{i,\alpha}^2$  de cet indice peuvent être calculés, avec un raisonnement semblable à celui de Box [5], par :

$$\tau_{i,\alpha}^2 = g^{(i)} \chi_{h^{(i)},\alpha}^2 \quad (3.77)$$

où

$$g^{(i)} = \theta_2^{(i)} / \theta_1^{(i)} = \sum_{j=m-i+1}^m \lambda_j^2 / \sum_{j=m-i+1}^m \lambda_j \quad (3.78)$$

$$h^{(i)} = \theta_1^{2(i)} / \theta_2^{(i)} = \left( \sum_{j=m-i+1}^m \lambda_j \right)^2 / \sum_{j=m-i+1}^m \lambda_j^2 \quad (3.79)$$

Pour améliorer la détection, on utilise le filtre EWMA (3.63). Dans le cas de cet indice de détection, les dernières composantes filtrées sont données par l'équation suivante :

$$\bar{t}_j(k) = (1 - \gamma) \bar{t}_j(k-1) + \gamma t_j(k) \quad (3.80)$$

où  $\bar{t}_j(k)$  est la  $j^{eme}$  composante filtrée à l'instant  $k$ . Ainsi, l'indice de détection filtré sera donné par l'équation :

$$\bar{D}_i(k) = \sum_{j=m-i+1}^m \bar{t}_j^2(k) \quad i = 1, 2, \dots, (m - \ell) \quad (3.81)$$

Les seuils de détection des indices  $\bar{D}_i$  seront calculés comme dans le cas du  $\overline{SPE}$  (3.75), par :

$$\bar{\tau}_{i,\alpha}^2 = \frac{\gamma}{2 - \gamma} \tau_{i,\alpha}^2 \quad (3.82)$$

Il faut noter que, dans le cas où ces seuils de détection  $\tau_i$  ( $i = 1, 2, \dots, (m - \ell)$ ), correspondant aux différents indices  $D_i$ , ne sont pas adéquats, ils peuvent être déterminés par apprentissage pendant la phase d'identification (en absence de défauts).

Ainsi, la procédure de détection proposée peut se résumer par :

1. Appliquer l'ACP sur les données pour déterminer le modèle ACP, soit  $\hat{T} \in \mathbb{R}^\ell$ .
2. Calculer l'indice  $\bar{D}_i$  ( $i = 1, \dots, (m - \ell)$ ) (3.81) et déterminer les seuils de détection  $\bar{\tau}_i$  pour chaque indice.
3. Surveiller le processus en utilisant la procédure de détection basée sur les différents sous-espace résiduels :
  - (a) Pour chaque mesure à l'instant  $k$ , sélectionner le sous-espace résiduel pour  $i = 1$
  - (b) Calculer  $\bar{D}_i(k)$   
 Si  $\bar{D}_i(k) > \bar{\tau}_i^2$  aller à (c),  
 sinon  $i = i + 1$ , aller à (b),  
 répéter jusqu'à ce que  $i = m - \ell$   
 Fin de la procédure de détection.
  - (c) Procédure de localisation.

L'exemple choisi pour illustrer cette procédure de détection sera l'exemple 2. Un modèle ACP à trois composantes a été retenu. Un défaut affectant la variable  $x_3$  avec une amplitude de 20% de la plage de variation de cette variable a été simulé. Le défaut simulé est un biais introduit à partir de l'instant 300. La figure (fig 3.9) présente l'évolution du  $SPE$ , tandis que la figure (fig 3.10) présente l'évolution du  $\overline{SPE}$ . Le défaut apparaît nettement mieux sur le  $SPE$  filtré que sur le  $SPE$ . Comme nous l'avons déjà expliqué, à cause du taux de fausses alarmes élevé que présente le  $SPE$  filtré, une augmentation du seuil s'impose mais dans ce cas le défaut ne peut pas être détecté.

Par contre, en utilisant l'indice  $\bar{D}_i$ , le défaut est détecté sur l'indice  $\bar{D}_3$  comme le montre la figure (fig 3.11).

Un autre défaut a été simulé sur la variable  $x_9$  à partir de l'instant 350, avec une amplitude d'environ 22% de la plage de variation de cette variable. La figure (fig 3.12) présente l'évolution du  $\overline{SPE}$  en présence du défaut. A cause des erreurs de modélisation, la détection de ce défaut est impossible. En appliquant la méthode de détection proposée, le défaut est détecté avec l'indice  $\bar{D}_1$  (fig 3.13).

Une comparaison quantitative entre les indices  $SPE$  et  $D_i$  est effectuée en exploitant les conditions suffisantes de détection des deux indices. Dans le cas du  $SPE$ , l'amplitude minimale

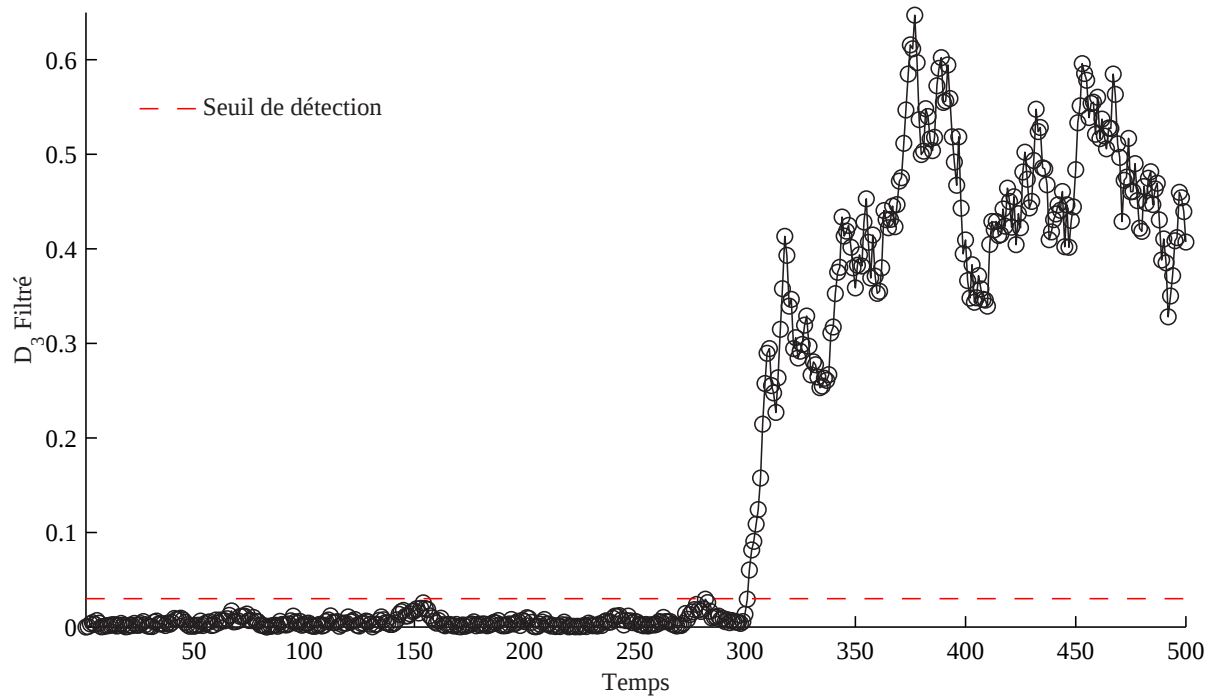


Figure 3.11 – Evolution de l'indice  $D_3$  filtré avec un défaut affectant la variable  $x_3$  à partir de l'instant 300

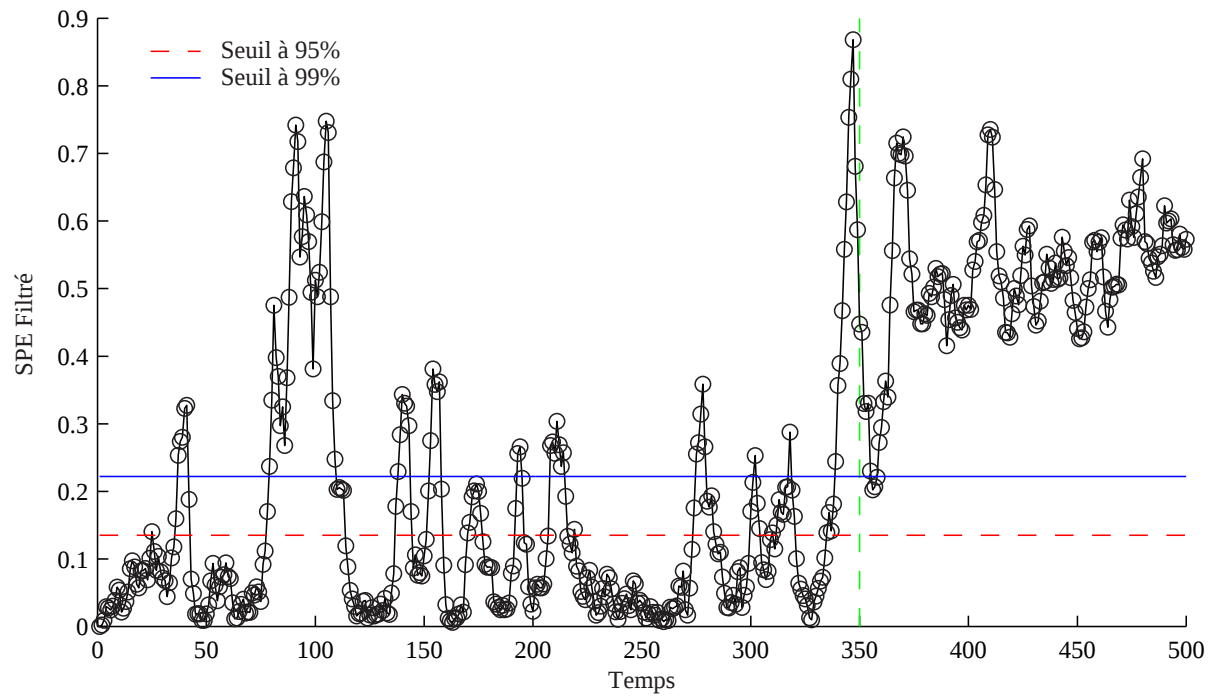


Figure 3.12 – Evolution du  $\overline{SPE}$  avec un défaut affectant la variable  $x_9$  à partir de l'instant 350

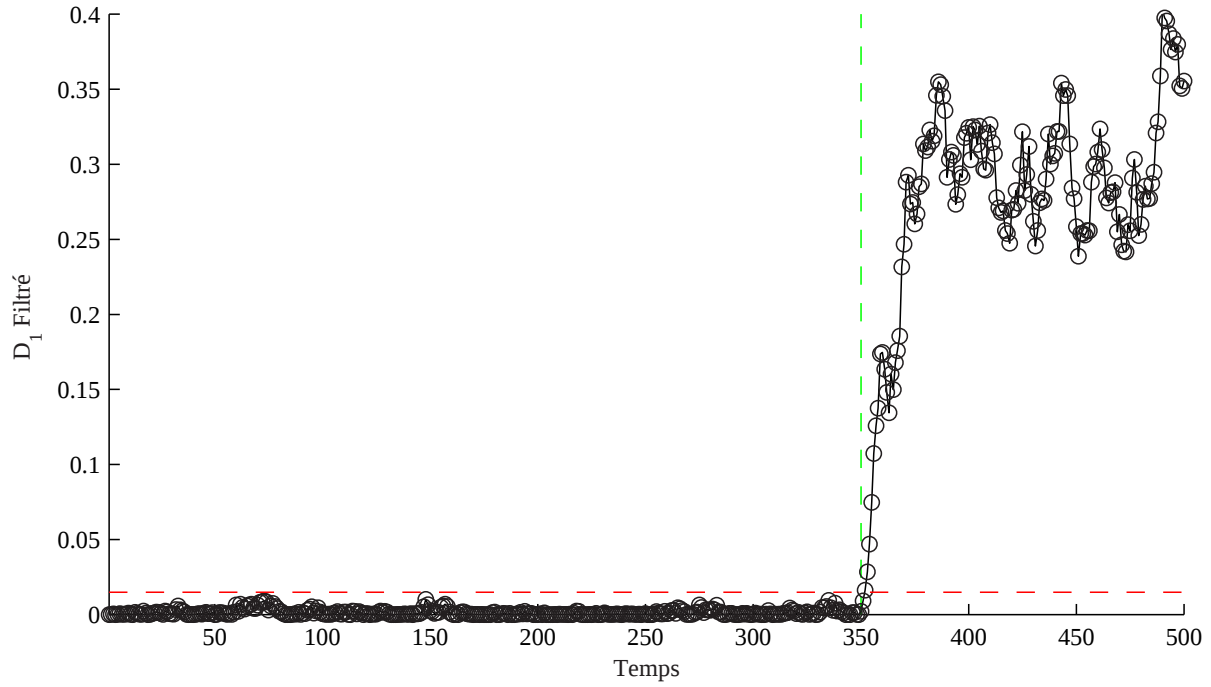


Figure 3.13 – Evolution de l'indice  $\bar{D}_1$  avec un défaut affectant la variable  $x_9$  à partir de l'instant 350

du défaut affectant la  $j^{eme}$  variable et vérifiant la condition suffisante de détectabilité est donnée par :

$$d_j = \frac{2\delta}{\|\tilde{\xi}_j\|} \quad (3.83)$$

Pour l'indice  $D_i$  cette condition sera donnée par l'équation :

$$d_j^{(i)} = \frac{2\tau_i}{\|\tilde{\xi}_j^{(i)}\|} \quad (3.84)$$

où  $d_j^{(i)}$  est l'amplitude du défaut affectant la  $j^{eme}$  variable et vérifiant la condition suffisante de détectabilité avec l'indice  $D_i$  ( $i = 1, \dots, (m - \ell)$ ).

$\xi_j^{(i)} = (I - \hat{C}_i)\xi_j$ ,  $\hat{C}_i = \hat{P}^{(i)}\hat{P}^{(i)T}$  et  $\hat{P}^{(i)}$  est formé par les  $i$  derniers vecteurs propres de la matrice  $\Sigma$ .

Le tableau (tab 3.2) présente les résultats du calcul des amplitudes minimales de défauts vérifiant les conditions suffisantes de détectabilité, dans le cas de l'exemple 1, pour les deux indices  $SPE$  et  $D_i$  respectivement.

Le tableau (tab 3.3) présente les résultats du calcul des amplitudes minimales de défauts vérifiant les conditions suffisantes de détectabilité, dans le cas de l'exemple 2, pour les deux indices  $SPE$  et  $D_i$  respectivement.

Il est clair que les amplitudes de défauts détectées par les indices  $D_i$  sont, dans la plupart des cas, nettement inférieures à celles détectées par le  $SPE$ .

|       | $d_1$ | $d_2$ | $d_3$ | $d_4$ | $d_5$ | $d_6$ | $d_7$ |
|-------|-------|-------|-------|-------|-------|-------|-------|
| $SPE$ | 0.80  | 0.80  | 0.80  | 0.92  | 0.92  | 0.80  | 0.79  |
| $D_1$ | 1.55  | 3.91  | 1.43  | 0.57  | 0.69  | 0.11  | 0.12  |
| $D_2$ | 0.55  | 0.52  | 0.62  | 0.62  | 0.60  | 0.19  | 0.19  |
| $D_3$ | 0.53  | 0.79  | 0.88  | 0.64  | 0.64  | 0.40  | 0.40  |
| $D_4$ | 0.69  | 0.64  | 0.98  | 0.81  | 0.82  | 0.60  | 0.60  |

Table 3.2 – Les amplitudes de défaut affectant les différentes variables et vérifiant la condition suffisante de détectabilité dans le cas du  $SPE$  et des indices  $D_i$  pour l'exemple 1.

|       | $d_1$ | $d_2$ | $d_3$ | $d_4$ | $d_5$ | $d_6$ | $d_7$ | $d_8$ | $d_9$ | $d_{10}$ |
|-------|-------|-------|-------|-------|-------|-------|-------|-------|-------|----------|
| $SPE$ | 5.87  | 5.85  | 5.85  | 6.03  | 5.99  | 6.93  | 5.88  | 7.62  | 5.77  | 6.11     |
| $D_1$ | 0.29  | 0.29  | 0.29  | 0.78  | 0.96  | 1.92  | 0.29  | 8.46  | 0.28  | 1.22     |
| $D_2$ | 0.52  | 0.53  | 0.52  | 0.83  | 1.07  | 3.16  | 0.52  | 3.21  | 0.51  | 0.54     |
| $D_3$ | 0.82  | 0.82  | 0.82  | 0.85  | 0.85  | 4.94  | 0.82  | 4.95  | 0.80  | 0.85     |
| $D_4$ | 1.22  | 1.24  | 1.22  | 3.27  | 4.02  | 8.05  | 1.22  | 35.44 | 1.19  | 5.12     |
| $D_5$ | 1.93  | 1.96  | 1.92  | 3.08  | 3.96  | 11.70 | 1.93  | 11.87 | 1.89  | 2.01     |
| $D_6$ | 2.57  | 2.58  | 2.57  | 2.66  | 2.66  | 15.48 | 2.59  | 15.52 | 2.53  | 2.68     |

Table 3.3 – Les amplitudes de défauts affectant les différentes variables et vérifiant la condition suffisante de détectabilité dans le cas du  $SPE$  et des indices  $D_i$  pour l'exemple 2.

Dans le deuxième exemple (tab 3.3), on constate que l'écart entre les amplitudes vérifiant la condition de détectabilité dans le cas de  $D_6$  et le  $SPE$  est beaucoup plus important que celui de l'exemple 1 (tab 3.2).

En analysant les expressions des deux indices, on constate que l'indice  $SPE$  est égal à l'indice  $D_6$  en ajoutant la composante  $t_4^2$ . Ainsi, on peut dire que dans le deuxième exemple on a plus d'incertitudes de modélisation et que la composante  $t_4$  est porteuse de ces incertitudes qui font que le  $SPE$  a une amplitude de défaut vérifiant la condition de détectabilité plus grande.

Si on définit la sensibilité  $\vartheta_i$  de l'indice  $D_i$  par rapport au  $SPE$  comme le rapport des amplitudes des défauts vérifiant la condition de détectabilité dans le cas de l'indice  $D_i$  et de l'indice  $SPE$  pour la même variable :

$$\vartheta_i = \frac{d_j}{d_j^{(i)}} \quad (3.85)$$

Ainsi on obtient le tableau suivant :

A partir de ce tableau il est clair que l'indice  $D_i$  est plus sensible aux défauts que le  $SPE$ . Par exemple l'indice  $D_3$  peut détecter un défaut affectant la variable  $x_3$  avec une amplitude 7 fois plus petite que celle que l'on peut détecter avec le  $SPE$ .

A partir de ces deux exemples, nous avons mis en évidence les limites de détection des indices  $SPE$  et  $\overline{SPE}$ . L'indice de détection proposé possède des capacités de détection supérieures à ceux-ci. De plus grâce à l'indice de détection proposé, des défauts de plus faible amplitude peuvent être détectés qu'avec le  $SPE$ , et ce du fait que les défauts affectant certaines variables sont détectables dans des sous-espaces résiduels de dimension bien inférieure à celle représentant le  $SPE$ .

|               | $d_1$ | $d_2$ | $d_3$ | $d_4$ | $d_5$ | $d_6$ | $d_7$ | $d_8$ | $d_9$ | $d_{10}$ |
|---------------|-------|-------|-------|-------|-------|-------|-------|-------|-------|----------|
| $\vartheta_1$ | 20.24 | 20.17 | 20.17 | 7.73  | 6.23  | 3.60  | 20.27 | 0.90  | 20.60 | 5.00     |
| $\vartheta_2$ | 11.28 | 11.03 | 11.25 | 7.26  | 5.59  | 2.19  | 11.30 | 2.37  | 11.31 | 11.31    |
| $\vartheta_3$ | 7.15  | 7.13  | 7.13  | 7.09  | 7.04  | 1.40  | 7.17  | 1.53  | 7.21  | 7.18     |
| $\vartheta_4$ | 4.81  | 4.71  | 4.79  | 1.84  | 1.49  | 0.86  | 4.81  | 0.21  | 4.84  | 1.19     |
| $\vartheta_5$ | 3.04  | 2.28  | 3.04  | 1.85  | 1.51  | 0.59  | 3.04  | 0.64  | 3.08  | 3.03     |
| $\vartheta_6$ | 2.28  | 2.26  | 2.27  | 2.26  | 2.25  | 0.44  | 2.27  | 0.49  | 2.28  | 2.27     |

Table 3.4 – Sensibilité des indices  $D_i$  aux défauts affectant les différentes variables par rapport au  $SPE$  pour l'exemple 2.

### 3.2.2 Génération de résidus par estimation paramétrique

L'autre approche de génération de résidus repose sur l'estimation des paramètres du modèle. Les résidus générés peuvent être directement l'écart entre les paramètres nominaux et les paramètres estimés. Dans le cadre de l'ACP, les résidus sont générés par un calcul d'angle entre les premières directions principales de référence (représentant le fonctionnement normal) et leur estimation.

L'idée de cette approche est d'effectuer une estimation en ligne des paramètres du modèle. Des résidus sont générés en comparant la valeur des paramètres en fonctionnement normal à l'estimation courante et le test des résidus obtenus détermine si les changements observés sont significatifs.

Dans le cas de l'ACP les paramètres du modèle sont estimés par décomposition en valeurs et vecteurs propres de la matrice de corrélation des mesures. La mise à jour des vecteurs propres nécessite une estimation de la matrice de corrélation après acquisition d'une nouvelle mesure ou d'un ensemble de mesures. Cette estimation de la matrice de corrélation peut s'effectuer soit par une mise à jour récursive (ACP récursive) [78] ou par une estimation de cette matrice à partir des échantillons d'une fenêtre glissante définie a priori [51]. Une autre solution pour estimer directement les vecteurs propres sans avoir besoin d'estimer la matrice de corrélation consiste à utiliser une approche neuronale [72, 84, 11].

L'indice de détection utilisé généralement avec cette approche est basé sur l'idée que le changement dans les conditions de fonctionnement se traduit par des changements des corrélations entre les variables du processus et peut être détecté par la surveillance des directions des composantes principales. La comparaison de ces vecteurs propres représentant le fonctionnement normal avec ceux calculés en ligne est réalisée par un indice de détection basé sur le calcul de l'angle entre les directions principales [80, 51]. L'indice proposé pour évaluer le changement des directions principales est donné par [51] :

$$In_i(k) = 1 - \left| p_i^T(k) p_{i0} \right| \quad i = 1, \dots, \ell \quad (3.86)$$

où  $p_i(k)$  est la  $i^{eme}$  direction principale calculée à l'instant  $k$  et  $p_{i0}$  représente la  $i^{eme}$  direction de référence,  $p_i^T(k) p_{i0}$  représente le cosinus de l'angle entre les deux vecteurs. Pour l'application de la méthode de détection proposée, les vecteurs  $p_{i0}$  de référence représentant le fonctionnement normal, les seuils de détection et les tailles des fenêtres glissantes utilisées pour le calcul des nouvelles directions  $p_i(k)$  doivent être déterminés. Ainsi, la procédure suivante est adoptée :

1. appliquer l'ACP pour modéliser le comportement du processus en fonctionnement normal sur un jeu d'identification et déterminer les directions  $p_{i0}$  de référence,

2. déterminer la taille de la fenêtre de surveillance.
3. A l'instant  $k$ , Appliquer l'ACP aux jeux de données de validation sélectionné par la fenêtre d'observation, et calculer les directions  $p_i(k)$  (fig 3.14),
4. calculer  $In_i(k)$ ,  $i = 1, \dots, \ell$ , mettre  $k = k + 1$ , glisser la fenêtre d'un pas et aller à l'étape 3.
5. déterminer les seuils des  $In_i$ .

Initialement le seuil de détection est déterminé à partir du jeux d'identification ensuite il est validé sur le jeu d'identification en choisissant une taille de fenêtre adéquate. Une fois le seuil de détection et la taille de fenêtre déterminés, la procédure de détection peut être appliquée en ligne.

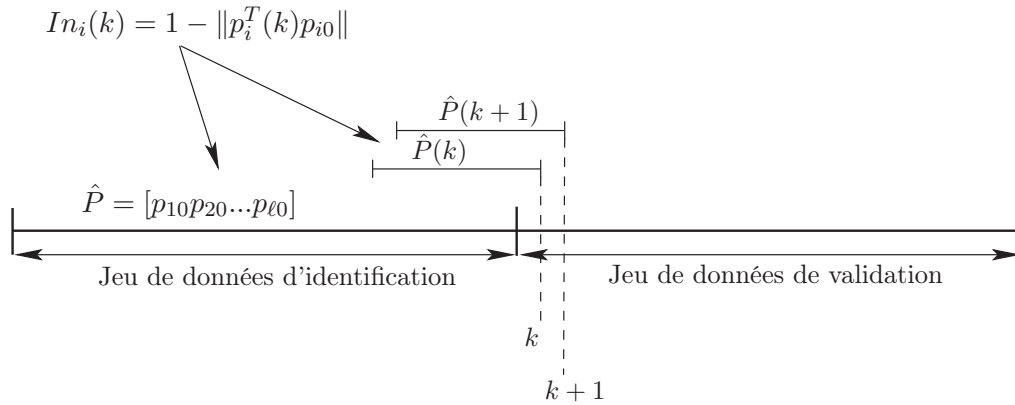


Figure 3.14 – Principe de la génération de résidus par estimation paramétrique dans le cas de l'ACP

Le seuil de détection est déterminé avec un certain taux de fausse alarme.

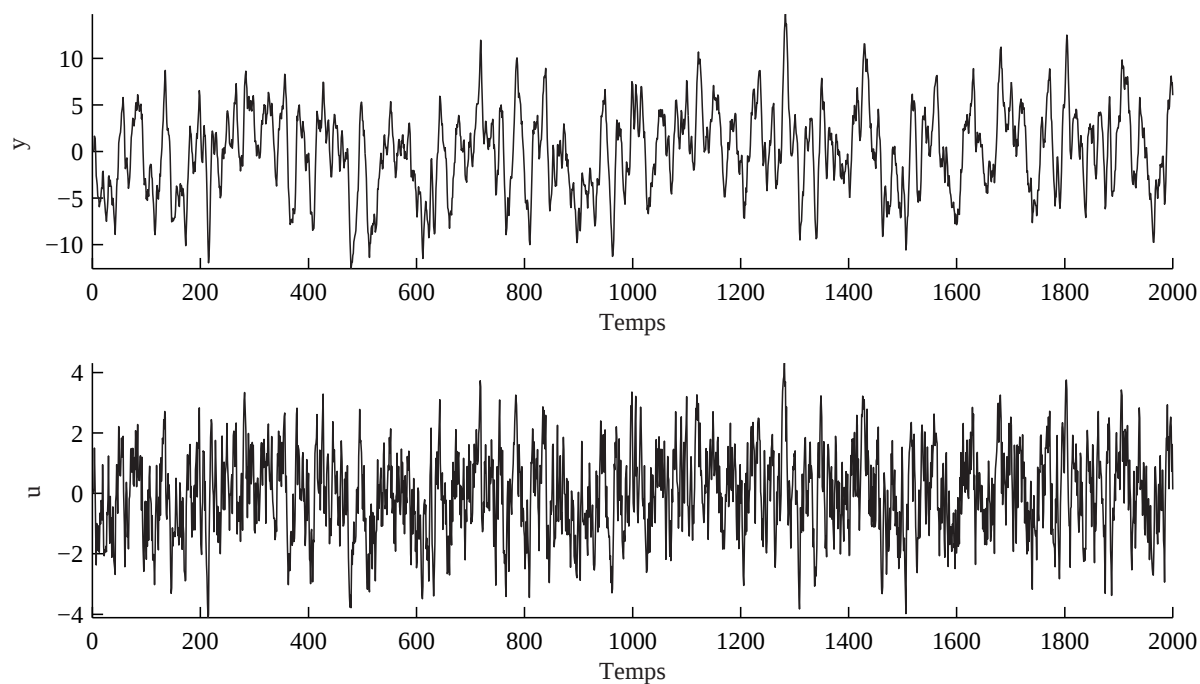
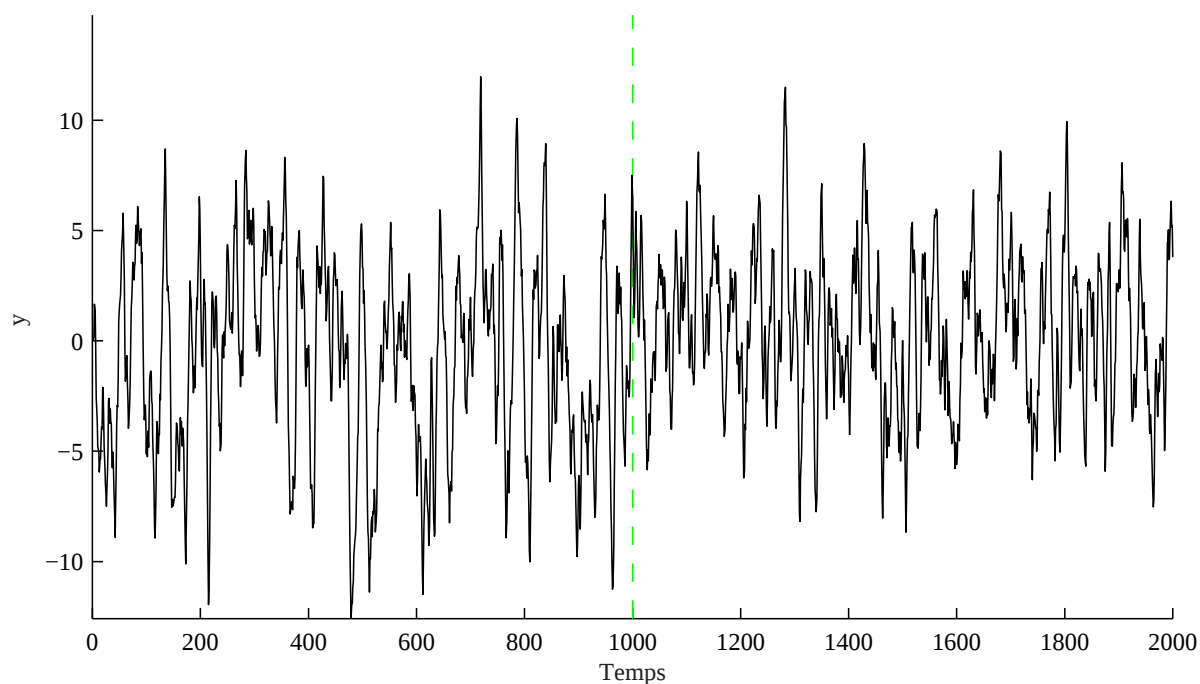
Il faut noter que la détermination de la taille de la fenêtre d'observation adéquate est très importante pour la qualité de la détection.

Comme les dernières directions principales sont plus sensibles aux bruits, la surveillance dans ce cas ne concerne que les premières directions principales.

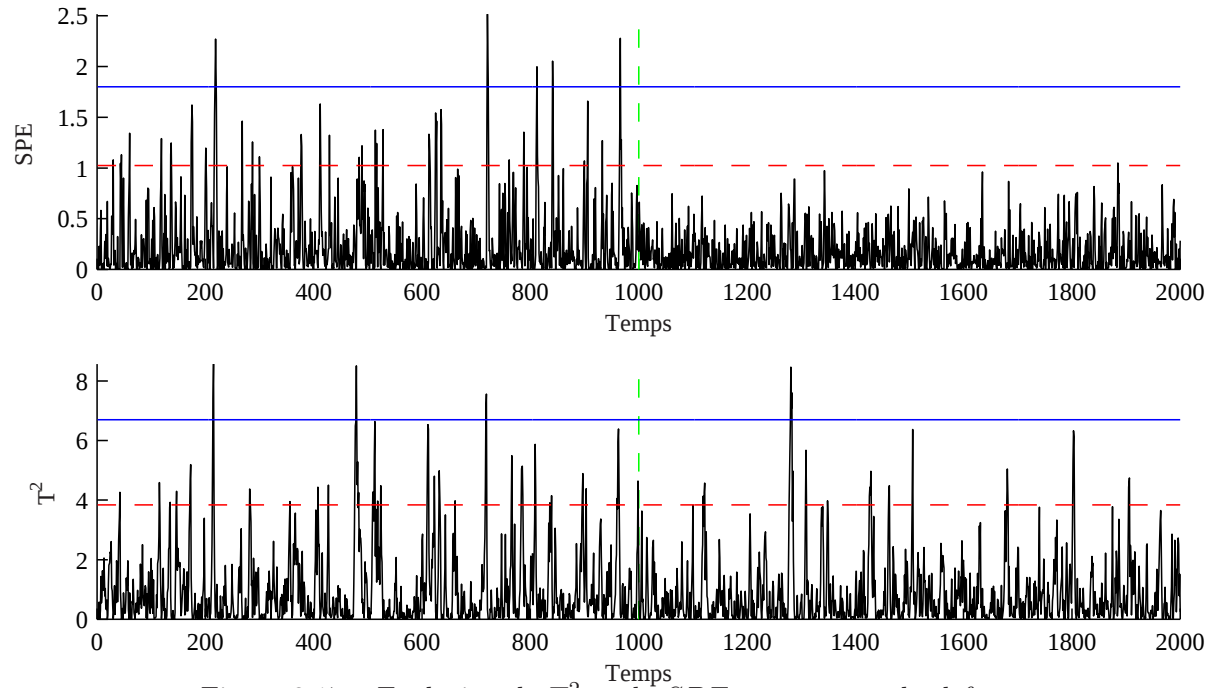
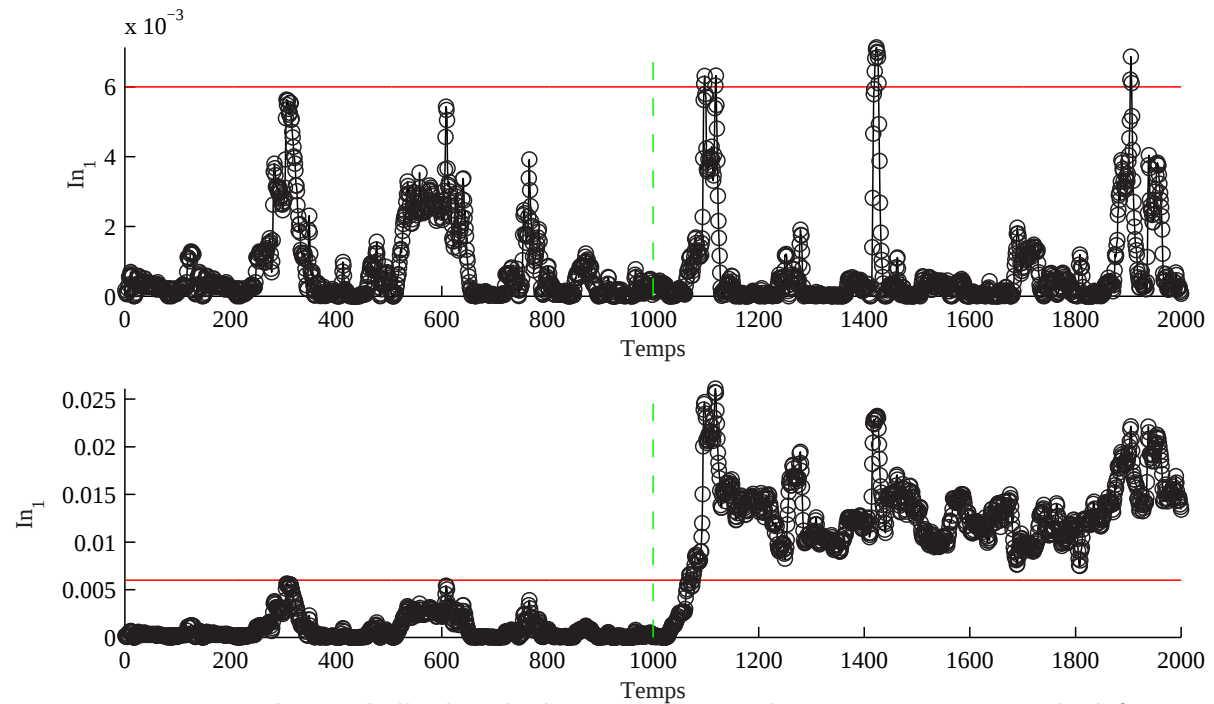
Pour illustrer cette méthode, un exemple simple de deux variables  $u$  et  $y$  représentant l'entrée et la sortie d'un système est présenté.

$$\begin{cases} y(k) = ay(k-1) + bu(k) \\ u(k) = cu(k-1) + dw(k) \end{cases} \quad (3.87)$$

où  $w$  est un bruit centré et  $a = 0.8$ ,  $b = 1$ ,  $c = 0.7$ ,  $d = 1$ . La figure (fig 3.15) présente l'évolution des deux variables  $y$  et  $u$  en fonctionnement normal, 2000 échantillons ont été générés. Un défaut d'environ 13% sur le paramètre  $a$  a été simulé, le nouveau paramètre est  $a = 0.7$  à partir de l'instant 1000. La figure (fig 3.16) présente l'évolution de la variable  $y$  en présence du défaut. Pour la détection du défaut, nous avons calculé dans un premier temps les deux indices  $T^2$  et  $SPE$  qui sont représentés sur la figure (fig 3.17), sachant qu'une seule composante a été retenue par le modèle ACP. On constate que le défaut simulé n'est pas détecté par les deux indices. La figure (fig 3.18) présente l'évolution de l'indice de détection  $In_1$  en absence et en présence du défaut, la taille de la fenêtre glissante utilisée est de 130 point. A partir de cet exemple simple nous avons montré les capacités de détection de l'indice  $In_i$  par rapport aux autres indices de détection.

Figure 3.15 – Evolution des deux variables  $y$  et  $u$  en absence de défautFigure 3.16 – Evolution de la variable  $y$  avec un changement de 13% sur le paramètre  $a$  à l'instant 1000.

Le principal inconvénient de cet indice est qu'il introduit un retard à la détection qui dépend directement de la taille de la fenêtre glissante utilisée.

Figure 3.17 – Evolution de  $T^2$  et de  $SPE$  en présence du défautFigure 3.18 – Evolution de l'indice de détection  $In_1$  en absence et en présence de défaut

### 3.3 Localisation de défauts

Lorsqu'un défaut est détecté, il est nécessaire d'identifier la ou les variables qui sont en cause : c'est la localisation de défauts. Pour réaliser cette tâche, plusieurs méthodes ont été développées dans la littérature. Concernant la localisation de défauts par l'analyse en composantes principales,

on retrouve trois approches :

- l’approche par structuration des résidus [26, 43, 79]. Cette approche s’inspire des méthodes de diagnostic utilisant la redondance analytique. L’évaluation de ces résidus permet de générer un vecteur de cohérence ou de signature de défauts représenté sous forme binaire suite à un test logique des résidus par rapport à leurs seuils respectifs [26]. Ces approches de localisation comparent la signature expérimentale et un ensemble de signatures théoriques également binaires. La localisation peut reposer soit sur le calcul de distance entre la signature expérimentale et les différentes signatures théoriques, soit sur une projection de la signature expérimentale sur les différentes signatures théoriques. Ainsi on cherche à transformer les résidus primaires en résidus secondaires ayant les propriétés de localisation recherchées.
- les approches utilisant des bancs de modèles. Dans cette catégorie d’approches on trouve trois approches :
  - l’approche utilisant des bancs de modèles ACP avec des ensembles de variables réduits. L’idée utilisée dans cette approche est la même que dans le cas des structures GOS (Generalized Observer Scheme) utilisé avec des observateurs. L’objectif est de générer des indicateurs de défaut sensible à certaines variables et pas aux autres.
  - approche utilisant le principe de reconstruction [15, 31]. Dans cette approche chaque capteur est supposé en défaut et reconstruit en utilisant le modèle ACP déjà calculé et les mesures des autres capteurs. La localisation est effectuée par comparaison de l’indice de détection avant et après reconstruction. La variable pour laquelle l’indice de détection après reconstruction est inférieur au seuil de détection, est déclarée en défaut.
  - l’approche par élimination est similaire à l’approche par reconstruction. Un ensemble d’indices de détection est généré en éliminant à chaque fois une variable de l’ensemble des variables à surveiller et en modifiant la matrice  $\hat{C}$  représentant le modèle ACP par élimination de la colonne correspondant à la variable éliminée. Le rapport entre la valeur des indices à l’instant de détection et leurs seuils est calculé. Le rapport calculé après l’élimination d’une variable ayant une valeur inférieure à 1 indique que cette variable est la variable incriminée.
- l’approche reposant sur le principe de calcul des contributions à l’indice de détection. Cette méthode consiste à calculer les contributions des différentes variables à l’indice de détection [62, 113, 68, 47, 108, 35], la variable ayant la plus grande contribution à la statistique de détection (calculée à l’instant de détection) est la variable incriminée.

### 3.3.1 Localisation par structuration des résidus

La méthode de localisation basée sur la structuration des résidus obtenus par ACP, a été présentée la première fois par Gertler *et al.* [25]. Cette approche consiste à chercher une transformation  $W$  de telle sorte que chaque résidu transformé soit sensible à certains défauts et insensible à d’autres ; le but est d’obtenir, pour chaque défaut une signature théorique permettant de localiser la variable en défaut. Une autre approche de structuration des résidus proposée par Huang *et al.* [43] consiste à utiliser des ACP partielles (ACP avec un nombre réduit de variables) et ainsi l’indice de détection ( $SPE$  dans son cas) obtenu pour chaque modèle partiel ne sera sensible qu’à un certain groupe de variables et pas aux autres. Dans le cas où un défaut affecte l’une des variables, une signature expérimentale est générée par évaluation des différents  $SPE$  correspondant aux différents modèles partiels permettant de déterminer la variable en défaut. Qin *et al.* [79] ont proposé une approche de structuration des résidus avec maximisation de la sensibilité aux défauts en se basant sur les travaux de Gertler *et al.* [25]. Dans cette section, nous allons

présenter le principe de ces approches de localisation.

### 3.3.1.1 Structuration de résidus

A partir des équations (3.25) et (3.26), un ensemble de résidus primaires est obtenu en présence de défaut affectant la  $j^{eme}$  variable :

$$\epsilon(k) = B\mathbf{x}(k) = B(\mathbf{x}^*(k) + \xi_j d(k)) = \epsilon^*(k) + B\xi_j d(k) \quad (3.88)$$

Pour améliorer les propriétés de localisation des résidus, ils doivent être transformés en :

$$\mathbf{r}(k) = W\epsilon(k) = W\epsilon^*(k) + WB\xi_j d(k) \quad (3.89)$$

où  $W$  est la matrice de transformation et  $\mathbf{r}(k)$  est le vecteur des résidus structurés. Les résidus structurés sont conçus de façon à ce que chaque résidu soit sensible à un sous-ensemble particulier de défauts. Une matrice d'incidence est formée, elle est constituée de 0 et de 1. Les lignes de cette matrice représentent la structure des résidus et les colonnes représentent les signatures des défauts. Chaque élément de cette matrice correspond à l'intersection d'un résidu  $r_i$  et d'un défaut  $d_j$ . Un 1 dans cette intersection signifie que le  $i^{eme}$  résidu est sensible au  $j^{eme}$  défaut alors qu'un 0 indique que ce résidu n'est pas sensible à ce défaut. Une fois la matrice des signatures théoriques formée, les lignes  $w_i^T$  de la matrice  $W$  sont choisies de façon à ce qu'on retrouve après multiplication par  $B$  les zéros correspondant sur la  $i^{eme}$  ligne de la matrice d'incidence :

$$w_i^T B^i = 0 \quad (3.90)$$

où  $B^i$  représente une matrice regroupant les colonnes de la matrice  $B$  correspondant à des zéros dans la  $i^{eme}$  ligne de la matrice d'incidence.

|       | $d_1$ | $d_2$ | $d_3$ | $d_4$ |
|-------|-------|-------|-------|-------|
| $r_1$ | 0     | 1     | 1     | 0     |
| $r_2$ | 0     | 0     | 1     | 1     |
| $r_3$ | 1     | 0     | 0     | 1     |
| $r_4$ | 1     | 1     | 0     | 0     |

Table 3.5 – Exemple d'une matrice de signature théorique dans le cas de quatre résidus

La recherche de la matrice  $W$  pour la structuration des résidus de l'exemple du tableau (tab 3.5) est effectuée de la manière suivante :

la première ligne de la matrice  $W$  doit vérifier la condition suivante :  $w_1^T [b_1 \ b_4] = [0 \ 0]$  ; la deuxième ligne doit satisfaire la condition :  $w_2^T [b_1 \ b_2] = [0 \ 0]$  et pour les deux autres lignes de cette matrice :  $w_3^T [b_2 \ b_3] = [0 \ 0]$  et  $w_4^T [b_3 \ b_4] = [0 \ 0]$ . Les  $b_i$  représentent les colonnes de la matrice  $B$ .

### Conditions d'existence de la transformation

Si nous avons  $m$  résidus primaires, alors les vecteurs  $w_i^T$  contiennent  $m$  éléments. Donc, le nombre de zéros dans chaque ligne de la matrice d'incidence ne doit pas dépasser  $(m - 1)$ . Plus précisément la matrice  $B^i$  doit vérifier la condition suivante :

$$\text{rang}(B^i) \leq m - 1 \quad (3.91)$$

De plus, pour ne pas avoir un zéro non voulu dans l'une des lignes de la matrice de structuration, c'est-à-dire pour que le résidu ne soit pas découplé d'un défaut qui apparaît avec un 1 sur l'une des lignes de la matrice de signatures théoriques, la condition :

$$\text{rang}[B^i \quad b_j] = \text{rang}(B^i) + 1 \quad (3.92)$$

doit être vérifiée pour toutes les colonnes  $b_j$  de la matrice  $B$  qui n'appartiennent pas à  $B^i$ .

### 3.3.1.2 Structuration des résidus avec l'ACP

Dans le cas de l'ACP, le vecteur des mesures, en fonctionnement normal, est donné par l'équation suivante :

$$\mathbf{x}^*(k) = \hat{P}\hat{\mathbf{t}}^*(k) + \tilde{P}\tilde{\mathbf{t}}^*(k) \quad (3.93)$$

Puisque  $\hat{P}$  et  $\tilde{P}$  sont orthogonaux, en pré-multipliant cette équation par  $\tilde{P}^T$ , nous obtenons le vecteur des résidus suivant, qui représente le vecteur des dernières composantes principales :

$$\tilde{P}^T \mathbf{x}^*(k) = \tilde{P}^T \tilde{P} \tilde{\mathbf{t}}^*(k) = \tilde{\mathbf{t}}^*(k) \quad (3.94)$$

Dans ce cas, en comparant les deux équations (3.26) et (3.94), nous obtenons :

$$B = \tilde{P}^T \quad (3.95)$$

Ainsi, l'équation (3.89) peut s'écrire dans ce cas sous la forme :

$$\mathbf{r}(k) = W\tilde{\mathbf{t}}(k) = W(\tilde{\mathbf{t}}^*(k) + \tilde{P}^T d(k)) \quad (3.96)$$

Donc, on doit chercher une matrice  $W$  telle que pour obtenir un zéro à l'intersection de la  $i^{eme}$  ligne et la  $j^{eme}$  colonne de la matrice de signatures théoriques, on doit avoir :  $w_i^T p_j = 0$  où  $p_j$  représente la  $j^{eme}$  colonne de la matrice  $\tilde{P}^T$ . Soit l'exemple 1 précédent, avec  $m = 7$  et  $\ell = 2$ . La matrice de signatures théoriques (d'incidence) est donnée par (tab 3.6) :

|       | $d_1$ | $d_2$ | $d_3$ | $d_4$ | $d_5$ | $d_6$ | $d_7$ |
|-------|-------|-------|-------|-------|-------|-------|-------|
| $r_1$ | 1     | 0     | 0     | 0     | 1     | 1     | 1     |
| $r_2$ | 1     | 1     | 0     | 0     | 0     | 1     | 1     |
| $r_3$ | 1     | 1     | 1     | 0     | 0     | 0     | 1     |
| $r_4$ | 1     | 1     | 1     | 1     | 0     | 0     | 0     |
| $r_5$ | 0     | 1     | 1     | 1     | 1     | 0     | 0     |
| $r_6$ | 0     | 0     | 1     | 1     | 1     | 1     | 0     |
| $r_7$ | 0     | 0     | 0     | 1     | 1     | 1     | 1     |

Table 3.6 – Matrice de signatures théoriques pour l'exemple 1

La matrice  $\tilde{P}^T$  est donnée par :

$$\tilde{P}^T = \begin{pmatrix} 0.41 & 0.24 & -0.67 & 0.38 & -0.40 & -0.00 & -0.01 \\ -0.36 & 0.70 & -0.36 & -0.35 & 0.33 & 0.00 & 0.03 \\ 0.60 & -0.30 & -0.30 & -0.48 & 0.47 & 0.00 & 0.00 \\ -0.31 & -0.32 & -0.27 & -0.26 & -0.27 & 0.45 & 0.60 \\ 0.02 & -0.01 & 0.02 & 0.06 & 0.05 & -0.75 & 0.65 \end{pmatrix}$$

Le calcul de la matrice de transformation  $W$  par cette méthode donne :

$$W = \begin{pmatrix} -0.11 & 0.14 & -0.57 & 0.75 & 0.26 \\ -0.12 & 0.71 & -0.61 & 0.06 & 0.31 \\ -0.83 & -0.32 & -0.45 & 0.03 & 0.01 \\ -0.82 & -0.32 & -0.45 & 0.03 & 0.01 \\ -0.85 & -0.26 & 0.44 & 0.02 & 0.01 \\ 0.69 & -0.25 & -0.42 & 0.37 & -0.35 \\ 0.55 & -0.48 & -0.67 & 0.00 & 0.02 \end{pmatrix}$$

Ainsi, le vecteur des résidus structurés est donnée par :

$$\mathbf{r}(k) = W\tilde{P}^T \mathbf{x}(k) \quad (3.97)$$

$$\mathbf{r}(k) = \begin{pmatrix} -0.67 & 0 & 0 & 0 & -0.37 & 0.13 & 0.62 \\ -0.69 & 0.63 & 0 & 0 & 0 & -0.21 & 0.26 \\ -0.50 & -0.30 & 0.80 & 0 & 0 & 0 & 0.00 \\ -0.50 & -0.30 & 0.80 & 0.00 & 0 & 0 & 0 \\ 0 & -0.53 & 0.53 & -0.45 & 0.45 & 0 & 0 \\ 0 & 0 & -0.36 & 0.43 & -0.69 & 0.44 & 0 \\ 0 & 0 & 0 & 0.70 & -0.70 & -0.01 & 0.00 \end{pmatrix} \mathbf{x}(k) \quad (3.98)$$

Pour illustrer cette méthode, des défauts ont été simulés sur les différentes variables (un défaut simple à la fois) d'une amplitude d'environ 20% de la plage de variation de la variable considérée et les résidus obtenus sont représentés sur la figure (fig 3.19). A partir de cette figure, on constate de nombreuses fausses alarmes sur les résidus. De plus, les signatures expérimentales de certains défauts sont dégradées (exemple des défauts  $d_4$ ,  $d_6$  et  $d_7$ ).

Pour améliorer la qualité des résidus structurés, Qin *et al.* [79] proposent de maximiser la sensibilité de ces derniers aux défauts.

### 3.3.1.3 Résidus structurés avec maximisation de la sensibilité aux défauts

Dans le cas simple où chaque résidu doit être insensible à un seul défaut, soit le  $i^{eme}$  capteur, l'équation (3.94) s'écrit :

$$\tilde{\mathbf{t}}(k) = \tilde{\mathbf{t}}^*(k) + b_i d_i(k) \quad (3.99)$$

où  $b_i$  est la  $i^{eme}$  colonne de la matrice  $B = \tilde{P}^T$ . Dans la suite on utilisera la notation  $B$  au lieu de  $\tilde{P}^T$  pour simplifier.

L'ensemble des résidus structurés  $\mathbf{r}(k)$  peut être généré par multiplication du vecteur des résidus  $\tilde{\mathbf{t}}(k)$  par la matrice de transformation  $W$ . La matrice  $W$  est conçue de telle sorte que chaque élément de  $\mathbf{r}(k)$  soit insensible à un défaut capteur particulier et sensible aux autres défauts. Le  $i^{eme}$  élément de  $\mathbf{r}(k)$  est donné par l'équation suivante :

$$r_i(k) = w_i^T \tilde{\mathbf{t}}(k) = w_i^T \tilde{\mathbf{t}}^*(k) + w_i^T b_i d_i(k) \quad (3.100)$$

avec  $i = 1, 2, \dots, m$  et où  $w_i^T$  est la  $i^{eme}$  ligne de la matrice  $W \in \mathbb{R}^{m \times \ell}$ .

Puisque  $\tilde{\mathbf{t}}^*(k)$  est un vecteur statistiquement proche de zéro représentant l'erreur d'estimation en l'absence de défaut, on peut générer le résidu  $r_i(k)$  pour qu'il soit insensible à  $d_i$ . Si le défaut actuel est  $d_j$  ( $j \neq i$ ), le résidu est donné par :

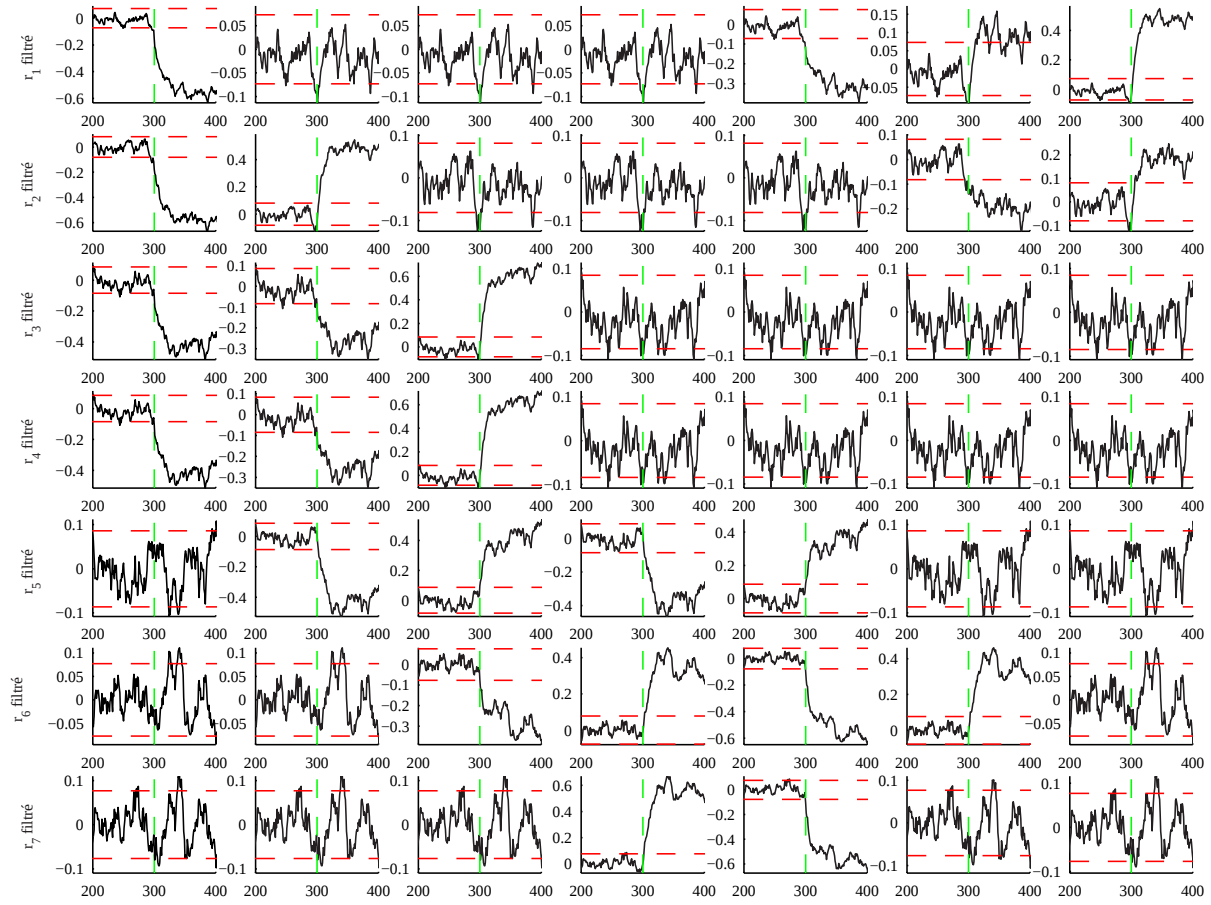


Figure 3.19 – Evolution des différents résidus filtrés des défauts affectant les différentes variables (exemple 1)

$$r_i^j(k) = w_i^T \tilde{\mathbf{t}}^*(k) + w_i^T b_j d_j(k) \quad (3.101)$$

On cherche un vecteur  $w_i$  telle que  $r_i(k)$  soit insensible à  $d_i$  et sensible le plus possible à  $d_j$  ( $j \neq i$ ). Choisir  $w_i$  tel que  $r_i(k)$  est insensible au capteur  $i$  mais plus sensible aux autres. Mathématiquement ceci est équivalent à :

$$\max_{w_i} \sum_{j \neq i} \frac{(w_i^T b_j)^2}{\|w_i\|^2 \|b_j\|^2} \quad (3.102)$$

sous la contrainte  $w_i^T b_i = 0$ .

Géométriquement,  $w_i$  est choisi orthogonal à  $b_i$  tout en minimisant les angles qu'il forment avec les autres directions de défaut  $b_j$  ( $j \neq i$ ).

Nous allons considérer le cas où chaque résidu structuré est insensible à plusieurs défauts à la fois. Supposons que le  $i^{ieme}$  résidu structuré doit être insensible à un groupe de capteurs avec un ensemble d'indices  $\mathbf{g}_i$  où  $i \in \mathbf{g}_i$ , et plus sensible aux autres capteurs avec l'ensemble d'indices  $\bar{\mathbf{g}}_i$ . Ainsi les colonnes de la matrice  $B$  sont regroupés en  $B_{\mathbf{g}_i}$  et  $B_{\bar{\mathbf{g}}_i}$  et l'équation (3.94) peut s'écrire comme :

$$\tilde{\mathbf{t}}(k) = B\mathbf{x}(k) = \begin{bmatrix} B_{\mathbf{g}_i} & B_{\bar{\mathbf{g}}_i} \end{bmatrix} \begin{bmatrix} \mathbf{x}_{\mathbf{g}_i}(k) \\ \mathbf{x}_{\bar{\mathbf{g}}_i}(k) \end{bmatrix} \quad (3.103)$$

Le  $i^{eme}$  résidu structuré ( $i \in \mathbf{g}_i$ ) est :

$$r_i(k) = w_i^T \tilde{\mathbf{t}}(k) = w_i^T B_{\mathbf{g}_i} \mathbf{x}_{\mathbf{g}_i}(k) + w_i^T B_{\bar{\mathbf{g}}_i} \mathbf{x}_{\bar{\mathbf{g}}_i}(k) \quad (3.104)$$

La structuration nécessite que  $r_i(k)$  soit insensible aux défauts dans  $\mathbf{x}_{\mathbf{g}_i}$  mais plus sensible aux défauts dans  $\mathbf{x}_{\bar{\mathbf{g}}_i}$ . Mathématiquement,  $w_i$  doit être choisi comme :

$$J = \max_{w_i} \|w_i^T B_{\bar{\mathbf{g}}_i}\|^2 \quad (3.105)$$

sous les contraintes  $w_i^T B_{\mathbf{g}_i} = 0$  et  $\|w_i\| = 1$ . On définit le Lagrangien suivant :

$$\mathcal{L} = \frac{1}{2} w_i^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i + \lambda^T B_{\mathbf{g}_i}^T w_i + \frac{1}{2} \mu (w_i^T w_i - 1) \quad (3.106)$$

$$\frac{\partial \mathcal{L}}{\partial w_i} = B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i + B_{\mathbf{g}_i} \lambda + \mu w_i = 0 \quad (3.107)$$

$$\frac{\partial \mathcal{L}}{\partial \lambda} = B_{\mathbf{g}_i}^T w_i = 0 \quad (3.108)$$

$$\frac{\partial \mathcal{L}}{\partial \mu} = \frac{1}{2} (w_i^T w_i - 1) = 0 \quad (3.109)$$

En multipliant l'équation (3.107) à gauche par  $w_i^T$  et en tenant compte de (3.109), on obtient :

$$w_i^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i + \mu = 0 \quad (3.110)$$

En multipliant l'équation (3.107) à gauche par  $B_{\mathbf{g}_i}^T$ , on obtient :

$$B_{\mathbf{g}_i}^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i + B_{\mathbf{g}_i}^T B_{\mathbf{g}_i} \lambda + \mu B_{\mathbf{g}_i}^T w_i = 0 \quad (3.111)$$

$$\lambda = - \left( B_{\mathbf{g}_i}^T B_{\mathbf{g}_i} \right)^{-1} B_{\mathbf{g}_i}^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i \quad (3.112)$$

Le report des équations (3.110) et (3.112) dans l'équation (3.107) nous donne :

$$B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i - B_{\mathbf{g}_i} \left( B_{\mathbf{g}_i}^T B_{\mathbf{g}_i} \right)^{-1} B_{\mathbf{g}_i}^T \left( B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T \right) w_i - \left( w_i^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i \right) w_i = 0 \quad (3.113)$$

$$\left( B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T \left( I - B_{\mathbf{g}_i} \left( B_{\mathbf{g}_i}^T B_{\mathbf{g}_i} \right)^{-1} B_{\mathbf{g}_i}^T \right) - \left( w_i^T B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T w_i \right) I \right) w_i = 0 \quad (3.114)$$

Posons :  $\mathcal{A} = B_{\bar{\mathbf{g}}_i} B_{\bar{\mathbf{g}}_i}^T$  et  $\mathcal{Q} = \left( I - B_{\mathbf{g}_i} \left( B_{\mathbf{g}_i}^T B_{\mathbf{g}_i} \right)^{-1} B_{\mathbf{g}_i}^T \right)$

$$\left( \mathcal{A} \mathcal{Q} - w_i^T \mathcal{A} w_i I \right) w_i = \left( \mathcal{A} \mathcal{Q} - \alpha I \right) w_i = 0 \quad (3.115)$$

Rappelons que l'on cherche  $w_i$  qui maximise le critère  $J$ , ainsi  $w_i$  est le vecteur propre de la matrice  $\mathcal{A} \mathcal{Q}$  correspondant à la plus grande valeur propre.

La différence entre la méthode de structuration avec maximisation de la sensibilité aux défauts et la méthode conventionnelle de structuration des résidus peut être illustrée géométriquement.

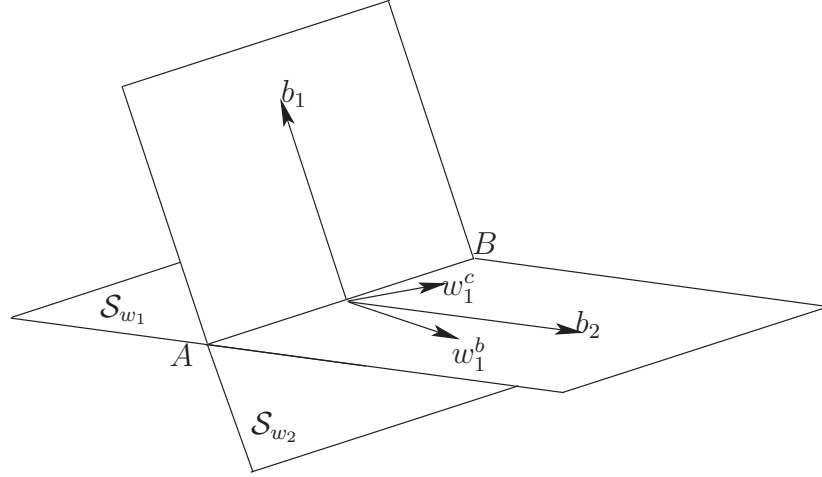


Figure 3.20 – Maximisation du produit  $w_1^b b_2$  et  $w_1^c$  est orthogonal à  $b_1$

La figure (fig 3.20) montre le cas de deux défauts capteurs  $b_1$  et  $b_2$  dans un espace tridimensionnel. Les plans orthogonaux à  $b_1$  et  $b_2$  sont notés par  $\mathcal{S}_{w1}$  et  $\mathcal{S}_{w2}$ , respectivement, qui passent par la même droite  $\overline{AB}$ . L'approche avec maximisation de la sensibilité aux défauts cherche à trouver  $w_1^b$  dans le plan  $\mathcal{S}_{w1}$  qui a l'angle le plus petit avec  $b_2$  (le plus sensible). L'approche conventionnelle de structuration de résidus peut choisir  $w_1^c$  arbitrairement dans le plan  $\mathcal{S}_{w1}$ . Un cas possible pour le choix de  $w_1^c$  est qu'il soit colinéaire avec  $\overline{AB}$  ce qui rend l'isolation des défauts  $b_1$  et  $b_2$  impossible.

Pour illustrer cette méthode, nous allons considérer l'exemple 1 avec la même matrice de signatures théoriques précédente (tab 3.6).

La matrice  $B = \tilde{P}^T$  contient cinq colonnes ( $m - \ell = 5$ ) et  $\mathbf{g}_i$  contient trois indices. Avec les conditions  $w_i^T B_{\mathbf{g}_i} = 0$  pour  $i \in \mathbf{g}_i$  et  $\|w_i\| = 1$ , la matrice de transformation  $W$  est donnée par :

$$W = \begin{pmatrix} 0.02 & -0.05 & -0.11 & -0.03 & -0.99 \\ 0.02 & 0.08 & -0.04 & -0.22 & -0.96 \\ -0.83 & -0.32 & -0.45 & 0.03 & 0.01 \\ 0.82 & 0.32 & 0.45 & -0.03 & -0.01 \\ -0.09 & 0.82 & 0.55 & -0.02 & -0.02 \\ 0.54 & -0.49 & -0.67 & -0.01 & 0.02 \\ -0.05 & 0.00 & 0.05 & -0.04 & -0.99 \end{pmatrix}$$

Ainsi le vecteur des résidus structuré résultant, est donné par l'expression suivante :

$$\mathbf{r}(k) = \begin{pmatrix} -0.07 & 0 & 0 & 0 & -0.14 & 0.76 & -0.62 \\ 0.00 & 0.16 & 0 & 0 & 0 & 0.62 & -0.76 \\ -0.50 & -0.30 & 0.80 & 0 & 0 & 0 & 0.00 \\ 0.50 & 0.30 & -0.80 & 0.00 & 0 & 0 & 0 \\ 0 & 0.39 & -0.39 & -0.58 & 0.58 & 0 & 0 \\ 0 & 0 & 0.01 & 0.71 & -0.70 & -0.02 & 0 \\ 0 & 0 & 0 & -0.10 & 0.00 & 0.72 & -0.67 \end{pmatrix} \mathbf{x}(k) \quad (3.116)$$

A partir de cette équation, on constate que cette approche de maximisation, pour cet exemple, n'a rien apporté de plus par rapport à l'approche classique. Certains coefficients de la matrice  $WB = W\tilde{P}^T$  sont améliorés alors que d'autres sont devenus plus petits qu'avant. On constate aussi que, par exemple, les deux défauts  $d_6$  et  $d_7$  ne sont pas localisables car ils ont la même signature expérimentale (fig 3.21).

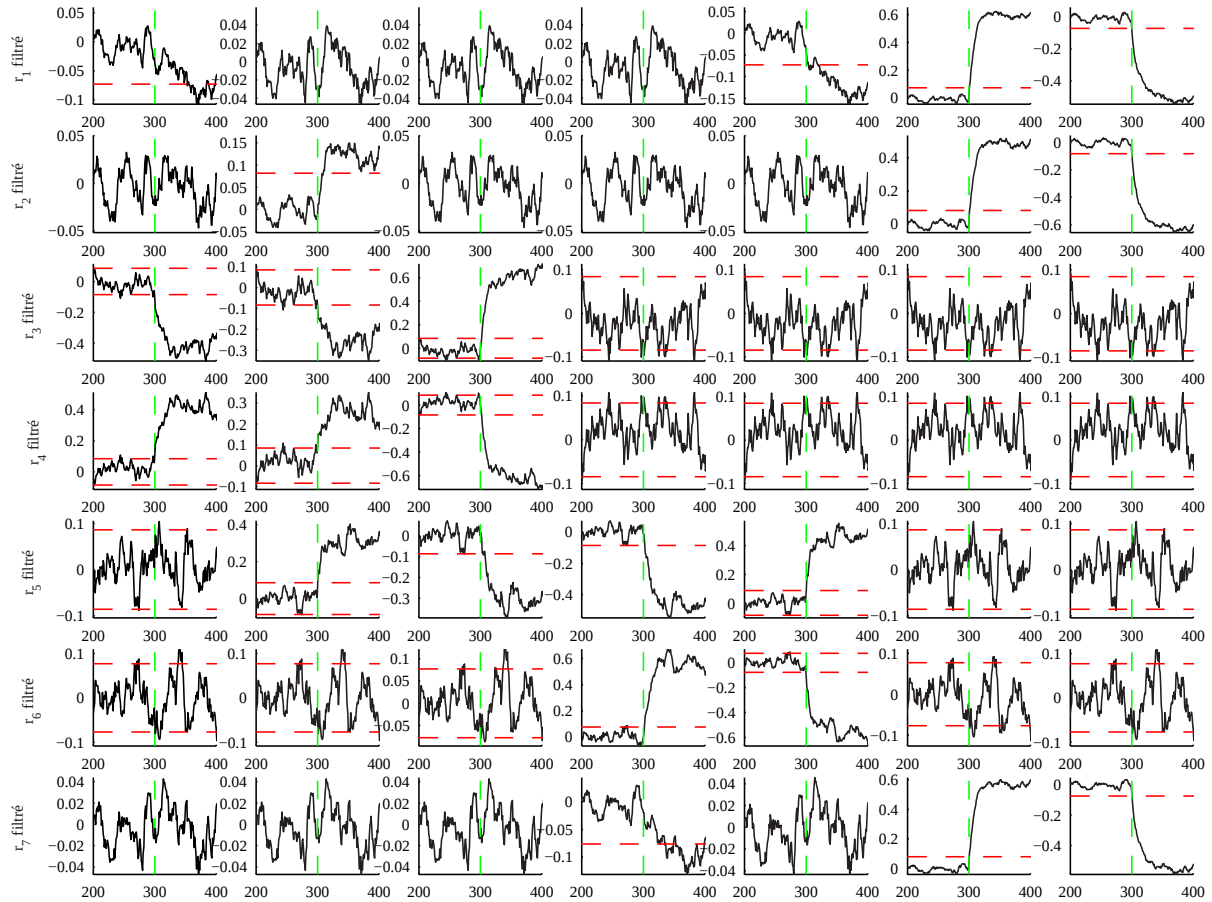


Figure 3.21 – Evolutions des différents résidus structurés filtrés pour des défauts affectant les différentes variables (exemple 1)

Considérons l'exemple 2. La matrice d'incidence, indiquant les occurrences des variables dans les différents résidus ainsi que l'influence des défauts sur les résidus, est donnée par :

|          | $d_1$ | $d_2$ | $d_3$ | $d_4$ | $d_5$ | $d_6$ | $d_7$ | $d_8$ | $d_9$ | $d_{10}$ |
|----------|-------|-------|-------|-------|-------|-------|-------|-------|-------|----------|
| $r_1$    | 1     | 0     | 0     | 0     | 0     | 0     | 1     | 1     | 1     | 1        |
| $r_2$    | 1     | 1     | 0     | 0     | 0     | 0     | 0     | 1     | 1     | 1        |
| $r_3$    | 1     | 1     | 1     | 0     | 0     | 0     | 0     | 0     | 1     | 1        |
| $r_4$    | 1     | 1     | 1     | 1     | 0     | 0     | 0     | 0     | 0     | 1        |
| $r_5$    | 1     | 1     | 1     | 1     | 1     | 0     | 0     | 0     | 0     | 0        |
| $r_6$    | 0     | 1     | 1     | 1     | 1     | 1     | 0     | 0     | 0     | 0        |
| $r_7$    | 0     | 0     | 1     | 1     | 1     | 1     | 1     | 0     | 0     | 0        |
| $r_8$    | 0     | 0     | 0     | 1     | 1     | 1     | 1     | 1     | 0     | 0        |
| $r_9$    | 0     | 0     | 0     | 0     | 1     | 1     | 1     | 1     | 1     | 0        |
| $r_{10}$ | 0     | 0     | 0     | 0     | 0     | 1     | 1     | 1     | 1     | 1        |

Table 3.7 – Matrice de signatures théoriques pour l'exemple 2

Le calcul de la matrice de transformation  $W$  par la méthode de structuration simple donne :

$$W = \begin{pmatrix} 0.12 & -0.02 & -0.02 & -0.16 & 0.17 & -0.81 & 0.50 \\ -0.05 & -0.09 & -0.46 & -0.62 & 0.26 & 0.24 & 0.50 \\ -0.04 & -0.18 & -0.21 & -0.94 & -0.02 & -0.08 & 0.10 \\ -0.04 & -0.24 & -0.24 & -0.92 & -0.04 & -0.17 & 0.00 \\ -0.03 & 0.85 & 0.10 & -0.06 & 0.48 & -0.13 & 0.04 \\ 0.04 & -0.82 & 0.00 & 0.53 & 0.18 & -0.01 & 0.03 \\ 0.04 & -0.85 & -0.09 & 0.09 & -0.39 & -0.23 & 0.19 \\ -0.98 & 0.00 & 0.06 & 0.056 & -0.06 & -0.14 & 0.06 \\ -0.97 & 0.10 & 0.06 & 0.03 & -0.06 & -0.14 & 0.06 \\ -0.95 & -0.04 & 0.22 & 0.03 & -0.08 & -0.10 & 0.15 \end{pmatrix} \quad (3.117)$$

et

$$\mathbf{r}(k) = \begin{pmatrix} -0.70 & 0 & 0 & 0 & 0 & 0 & 0.70 & -0.07 & 0.00 & 0.02 \\ -0.28 & 0.61 & 0 & 0 & 0 & 0 & 0 & -0.04 & -0.58 & 0.44 \\ -0.38 & 0.82 & -0.36 & 0 & 0 & 0 & 0 & 0 & -0.12 & 0.10 \\ -0.41 & 0.81 & -0.38 & -0.09 & 0 & 0 & 0 & 0 & 0 & 0.10 \\ -0.28 & -0.09 & 0.40 & 0.60 & -0.61 & 0 & 0 & 0 & 0 & 0 \\ 0 & -0.31 & 0.27 & -0.64 & 0.64 & -0.00 & 0 & 0 & 0 & 0 \\ 0 & 0 & -0.36 & -0.60 & 0.62 & -0.04 & 0.33 & 0 & 0 & 0 \\ 0 & 0 & 0 & -0.07 & -0.11 & -0.73 & 0.07 & 0.65 & 0 & 0 \\ 0 & 0 & 0 & 0 & -0.18 & -0.73 & 0.08 & 0.65 & 0.00 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.72 & 0.11 & 0.65 & -0.07 & -0.12 \end{pmatrix} \mathbf{x}(k) \quad (3.118)$$

De même, le calcul de la matrice de transformation  $W$  par la méthode de maximisation de la sensibilité aux défauts donne :

$$W = \begin{pmatrix} -0.12 & 0.02 & 0.02 & 0.16 & -0.17 & 0.81 & -0.50 \\ 0.00 & 0.15 & -0.25 & 0.61 & -0.32 & 0.58 & 0.30 \\ -0.04 & -0.18 & -0.21 & -0.94 & -0.02 & -0.08 & 0.10 \\ -0.04 & -0.24 & -0.24 & -0.92 & -0.04 & -0.17 & 0.00 \\ 0.01 & 0.13 & -0.02 & 0.27 & -0.87 & 0.37 & -0.10 \\ 0.01 & 0.17 & 0.13 & 0.82 & 0.52 & -0.02 & 0.05 \\ 0.02 & -0.23 & -0.07 & -0.10 & -0.74 & -0.47 & 0.38 \\ 0.98 & -0.00 & -0.06 & -0.05 & 0.06 & 0.14 & -0.06 \\ 0.97 & -0.10 & -0.07 & -0.03 & 0.06 & 0.14 & -0.06 \\ -0.93 & -0.03 & 0.13 & 0.04 & -0.07 & -0.07 & 0.30 \end{pmatrix} \quad (3.119)$$

et

$$\mathbf{r}(k) = \begin{pmatrix} 0.70 & 0 & 0 & 0 & 0 & 0 & -0.70 & 0.07 & 0.00 & -0.02 \\ 0.69 & -0.46 & 0 & 0 & 0 & 0 & 0 & -0.04 & -0.43 & 0.32 \\ -0.38 & 0.82 & -0.36 & 0 & 0 & 0 & 0 & 0 & -0.12 & 0.10 \\ -0.41 & 0.81 & -0.38 & -0.09 & 0 & 0 & 0 & 0 & 0 & 0.10 \\ 0.79 & -0.23 & -0.55 & 0.03 & -0.03 & 0 & 0 & 0 & 0 & 0 \\ 0 & -0.70 & 0.70 & 0.00 & -0.02 & -0.01 & 0 & 0 & 0 & 0 \\ 0 & 0 & -0.68 & -0.15 & 0.16 & -0.07 & 0.68 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0.07 & 0.11 & 0.73 & -0.07 & -0.65 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0.18 & 0.73 & -0.08 & -0.65 & 0.00 & 0 \\ 0 & 0 & 0 & 0 & 0 & -0.72 & 0.18 & 0.63 & -0.20 & -0.01 \end{pmatrix} \mathbf{x}(k) \quad (3.120)$$

La figure (fig 3.22) présente l'évolution des différents résidus structurés (par l'approche de maximisation de la sensibilité aux défauts) avec des défauts affectant les différentes variables.

A partir des deux équations (3.118) et (3.120), on constate que les deux défauts  $d_6$  et  $d_8$  ont quasiment la même influence sur les différents résidus, ce qui provoque une ambiguïté à localiser ces deux défauts (fig 3.22).

La mise en oeuvre de cette méthode de localisation est relativement aisée mais elle présente cependant un inconvénient majeur qui réside dans la fréquente dégradation des signatures expérimentales résultant de déclenchements excessifs ou partiels des tests. Le déclenchement excessif est provoqué par les perturbations et les erreurs de modélisation qui affectent les résidus (fig 3.22). Le déclenchement partiel des tests est lié au fait que la sensibilité des résidus aux défauts n'est pas prise en compte car la matrice d'incidence est choisie en fonction de la propriété de localisation (matrice fortement localisante) mais pas en fonction de la sensibilité des résidus aux défauts. Il est ainsi possible qu'un défaut ne soit pas détecté lorsque notamment celui-ci est de faible amplitude mais surtout lorsque la sensibilité du résidu à ce défaut est faible. La conséquence est alors une dégradation de la signature expérimentale. Comme l'illustre l'exemple précédent, une signature dégradée peut alors conduire à une mauvaise localisation, voire même à une impossibilité de localisation.

De plus, la condition d'existence (3.92) est très restrictive et rend la tâche de structuration non évidente. Il faut rappeler que cet exemple n'a pas pour but de remettre en question le principe de cette méthode, dont l'intérêt a été démontré à de nombreuses reprises dans la littérature, mais juste de mettre en évidence les limites de cette méthode.

### 3.3.2 Localisation utilisant un banc de modèles

La localisation utilisant un banc de modèles dans le cas de l'ACP, exploite le même principe utilisé avec des observateurs de type GOS (Generalized Observer Scheme) où le  $i^{eme}$  observateur

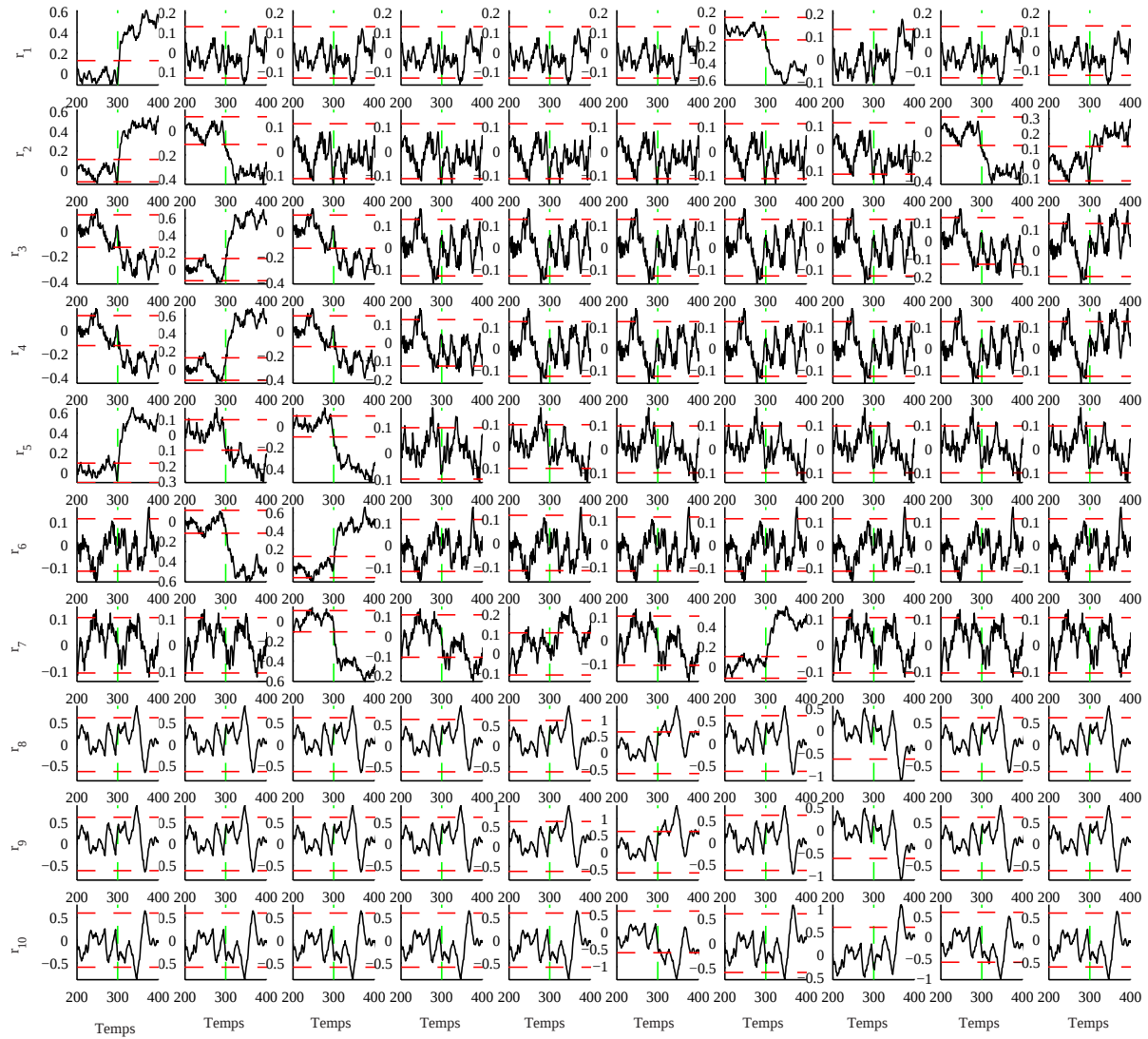


Figure 3.22 – Evolution des différents résidus structurés (filtrés avec un filtre EWMA,  $\gamma = 0.1$ ) pour des défauts affectant les différentes variables (exemple 2)

est piloté par toutes les entrées (sorties), sauf la  $i^{eme}$  et toutes les sorties (entrées). La sortie de cet observateur est donc sensible aux défauts de toutes les entrées (sorties) sauf ceux de la  $i^{eme}$  [20, 65]. Ainsi, on retrouve trois approches.

### 3.3.2.1 Localisation par ACP partielles

Une autre façon de structurer des résidus est d'utiliser ce que l'on appelle l'ACP partielle [43]. L'ACP offre une autre possibilité qui consiste à utiliser des modèles partiels. On sous entend par ACP partielle une ACP effectuée sur un vecteur de données réduit où quelques variables sont écartées par rapport au vecteur originel. Donc, les résidus vont être sensibles uniquement aux défauts associés aux variables qui forment le vecteur réduit. Les défauts associés aux variables éliminées n'affecteront pas les résidus correspondants. L'idée de l'ACP partielle a été décrite en premier par Gertler and Meters [25]. La procédure consiste à structurer les indices de détection en

calculant les ACP partielles ainsi que les seuils de détection des indices correspondants (fig 3.23).

#### Procédure de structuration des résidus

1. Appliquer l'ACP à la matrice des données.
2. Construire une matrice d'incidence fortement localisable (Matrice de signatures théoriques).
3. Construire un ensemble de modèles d'ACP partielles, chacune correspondant à une ligne de la matrice d'incidence (prendre les variables ayant un 1 sur cette ligne).
4. Déterminer les seuils pour la détection des défauts (seuil  $\tau_i^2$  sur l'indice  $SPE_i$ ).

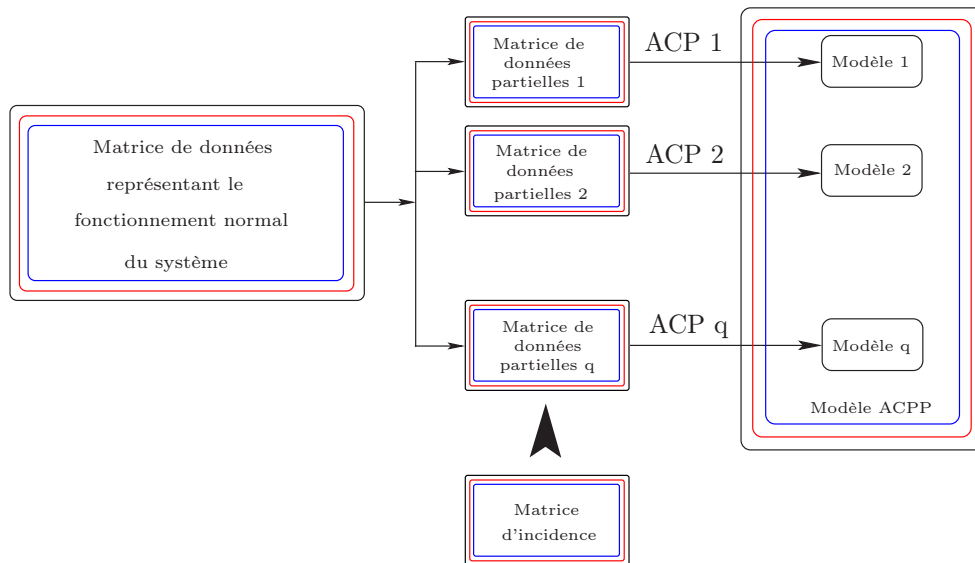


Figure 3.23 – Procédure de structuration de résidus par ACP partielles

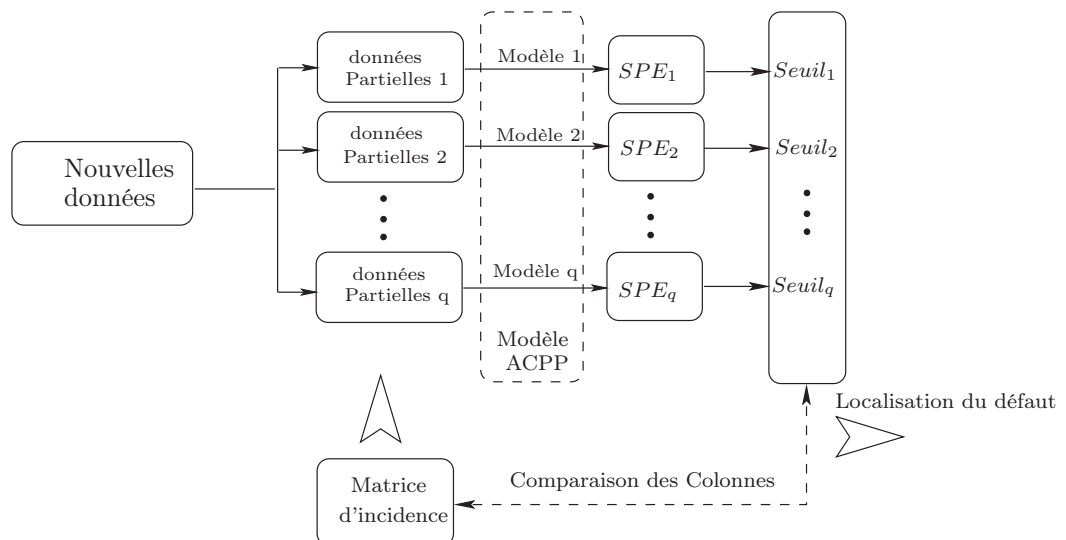


Figure 3.24 – Procédure de localisation par l'ACP partielle structurée

### Procédure de surveillance en ligne

Le test de localisation peut se faire en ligne, à chaque instant :

1. Calculer le  $SPE$  pour chacune des ACP partielles.
2. Comparer les indices aux seuils appropriés et former la signature expérimentale du défaut  $Se$  :

$$Se_i = \begin{cases} 0 & \text{si } SPE_i \leq Seuil_i \\ 1 & \text{si } SPE_i > Seuil_i \end{cases} \quad (3.121)$$

3. Comparer la signature expérimentale du défaut aux colonnes de la matrice d'incidence pour arriver à une décision de localisation.

Pour bien illustrer cette approche, considérons l'exemple 1. Ainsi et suivant la procédure de structuration par ACP partielles, un ensemble de sept ACP partielles ont été générées. La figure (fig 3.25) présente les réponses des indices de détection aux différents défauts. A partir de cette figure on constate que certains indices sont entachés d'erreurs de modélisation et qu'ils présentent des taux de fausses alarmes élevés. Cela peut s'expliquer par le manque de redondance dans les modèles ACP choisis. Considérons par exemple le quatrième modèle (le modèle à partir duquel de résidu  $SPE_4$  est généré) qui est calculé à partir des variables  $x_1$ ,  $x_2$ ,  $x_3$  et  $x_4$  (d'après la matrice de signatures théorique (tab 3.6)). Pour ce modèle une seule composante a été retenue et comme la variable  $x_4$  n'est pas corrélée avec les trois autres ( $x_1$ ,  $x_2$  et  $x_3$ ), elle va être projetée dans le sous-espace résiduel ce qui se traduit par des erreurs de modélisation sur l'indice de détection correspondant à ce modèle.

Le premier inconvénient de cette approche est qu'elle utilise l'indice  $SPE$  comme indicateur de défaut. De plus, à partir de cet exemple, on constate que le problème lié à cette méthode est que l'on ne sait pas s'il est possible d'élaborer un modèle ACP réduit a priori. Pour cette raison on impose une contrainte supplémentaire à cette approche qui consiste à s'assurer que pour chaque modèle partiel, toutes les variables doivent avoir, au moins, un degré de redondance égal à deux. Ainsi, on peut utiliser la variance non reconstruite pour la sélection des variables qui peuvent être utilisées dans chaque ACP partielle.

#### 3.3.2.2 Localisation par élimination

Stork *et al.* [89] a proposé une approche pour la localisation de défauts. La philosophie derrière cette méthode est similaire à celle de Dunia *et al.* [14]. Dans le cas où un défaut est détecté avec la statistique  $SPE$ , un ensemble de modèles ACP est généré à partir du modèle global en éliminant à chaque fois une variable de l'ensemble des variables à surveiller. Ainsi, la décomposition de la matrice de donnée  $X$  avec l'ACP, dont la  $j^{eme}$  variable  $X_j$  est éliminée, sera donnée par :

$$X_{-j} = T_{-j}P_{-j}^T + E_{-j} \quad (j = 1, 2, \dots, m) \quad (3.122)$$

La statistique  $SPE(k)$  est ainsi donnée, à l'instant  $k$ , par :

$$SPE_{-j}(k) = \mathbf{x}_{-j}^T(k) \left( I_{m-1} - \hat{P}_{-j}\hat{P}_{-j}^T \right) \mathbf{x}_{-j}(k) \quad (3.123)$$

où  $\mathbf{x}_{-j}(k)$  désigne le vecteur  $\mathbf{x}(k)$  à l'instant  $k$ , après élimination de la  $j^{ieme}$  variable,  $\hat{P}_{-j}$  représente la matrice  $\hat{P}$  du modèle ACP dont la  $j^{eme}$  ligne est éliminée. Cette procédure est répétée en éliminant une variable à chaque fois. La quantité définie par  $Q_r$  qui représente le rapport entre le  $SPE_{-j}$  et le seuil correspondant que l'on notera  $\delta_{\alpha, -j}^2$  [46] est ainsi calculée pour les  $m$  modèles :

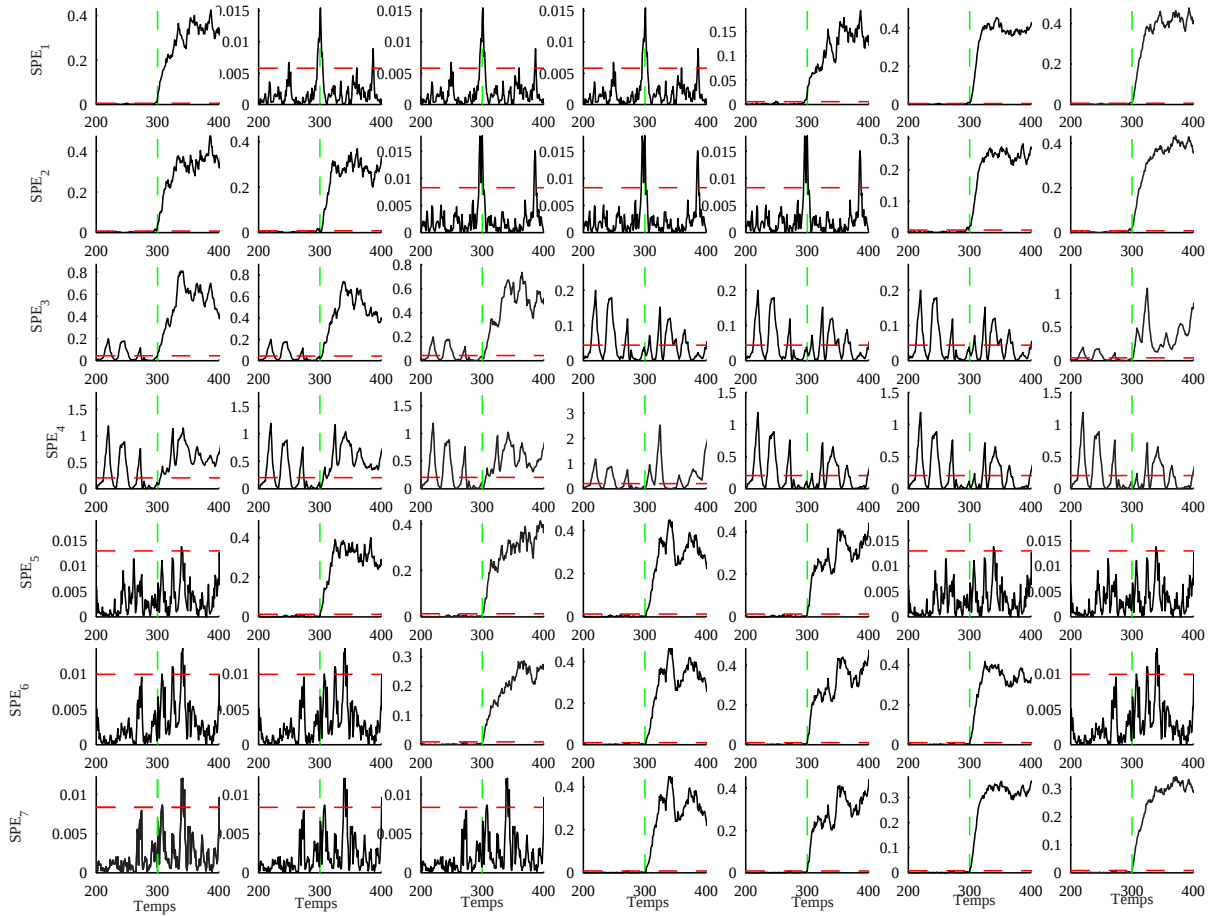


Figure 3.25 – Evolutions des différents  $SPE$  correspondant aux ACP partielles des différents défauts (exemple 1)

$$Q_r = \frac{SPE_{-j}}{\delta_{\alpha,-j}^2} \quad (3.124)$$

La variable éliminée, pour laquelle la valeur du rapport  $Q_r$  est la plus petite, est considérée comme la variable incriminée.

Cependant cette approche présente un inconvénient majeur que nous allons illustrer par l'exemple 1. La figure (fig 3.26) présente les différents résidus en absence de défaut. On constate que indices de détection  $SPE_i$  sont entachés d'erreur de modélisation et présentent un taux de fausses alarmes très élevé. En fait, l'élimination d'une ligne de la matrice  $\hat{P}$  crée un déséquilibre dans les équations permettant de générer les résidus (dans le cas où les équations sont équilibrées, les résidus générés sont statistiquement nuls).

Si, par contre, on augmente les seuils de détection de façon à ne plus être sensible à ces erreurs, on arrive à localiser la variable incriminée. Pour illustrer ce cas, un défaut sur la variable  $x_1$  d'une amplitude d'environ 20% de la plage de variation de cette variable, a été simulé entre les instants 350 et 500. Le résultat de l'application de l'algorithme d'élimination pour localiser la variable en défaut est représenté sur les figures (fig 3.27) et (fig 3.28).

Il faut noter que cet exemple simple (exemple 1) a été choisi pour illustrer le principe de la

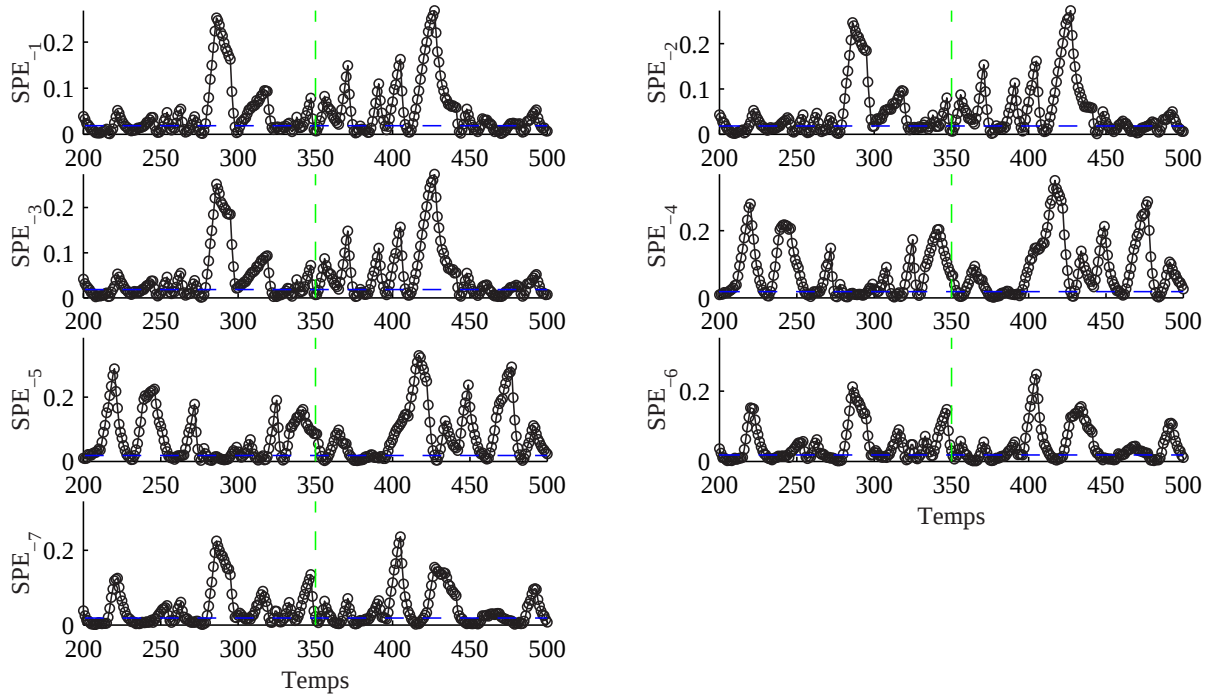


Figure 3.26 – Evolution des différents  $SPE$  obtenus par la méthode d'élimination en absence de défauts.

méthode. Comme le  $SPE$  est sensible aux erreurs de modélisation, il est souhaitable d'adapter cette approche à l'indice de détection proposé et qui a été décrit au début de ce chapitre.

Ainsi, reprenant l'exemple 2 avec un défaut affectant la variable  $x_3$ . Il faut rappeler ici que ce défaut a été détecté avec l'indice  $D_3$  (fig 3.11). L'application de la procédure proposée par Stork [89] sur l'indice  $D_3$  (utilisant un modèle à  $m - 3$  composantes au lieu de  $\ell$  composantes) permet de localiser la variable  $x_3$ . La figure (fig 3.29) présente l'évolution des différents indices  $D_3$  (qui ne sont en fait que des  $SPE$  calculer à partir d'un modèle à  $(m - 3)$  composantes). L'élimination de la variable  $x_3$  permet d'avoir un indice qui n'est pas affecté par le défaut, ce qui indique que la variable considérée est la variable incriminée.

### 3.3.2.3 Localisation basée sur le principe de reconstruction

Cette méthode basée sur le principe de reconstruction [14] consiste à suspecter qu'un capteur est défaillant et à reconstruire la valeur de sa mesure en se basant sur le modèle ACP déjà calculé et les mesures des autres capteurs (2.41). La procédure est répétée pour chaque capteur. La localisation est effectuée par comparaison de l'indice de détection avant et après reconstruction.

Dans un premier temps, nous allons rappeler brièvement le principe de la reconstruction ainsi que les différents résidus qui peuvent être obtenus à partir d'un modèle ACP.

Le vecteur de mesures reconstruit peut être utilisé pour définir un ensemble de résidus. Pour obtenir ces différents résidus, Dunia *et al.* [14] ont défini la matrice  $G_j$  :

$$G_j^T = [\xi_1 \ \xi_2 \ \dots \ g_j \ \dots \ \xi_m] \quad (3.125)$$

$$\text{où } g_j^T = \frac{1}{1-c_{jj}} [c_{-j}^T \ 0 \ c_{+j}^T].$$

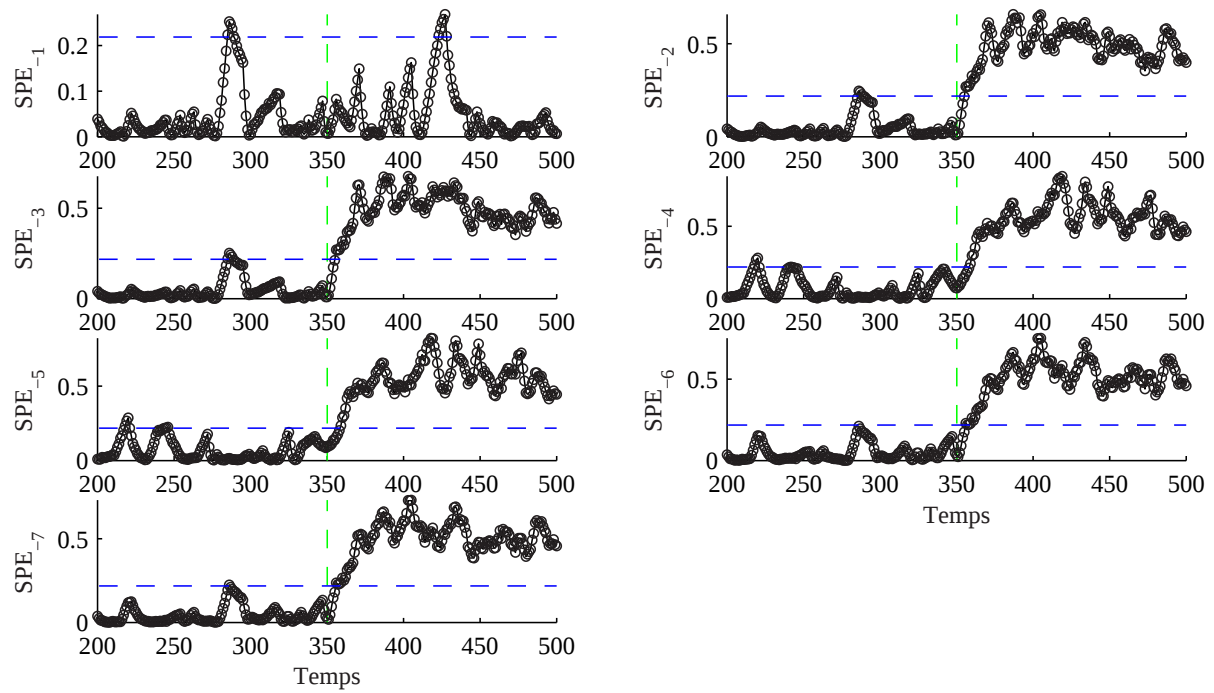


Figure 3.27 – Evolution des différents  $SPE$  obtenus par la méthode d'élimination avec un défaut affectant la variable  $x_1$  à partir de l'instant 350 avec seuils modifiés

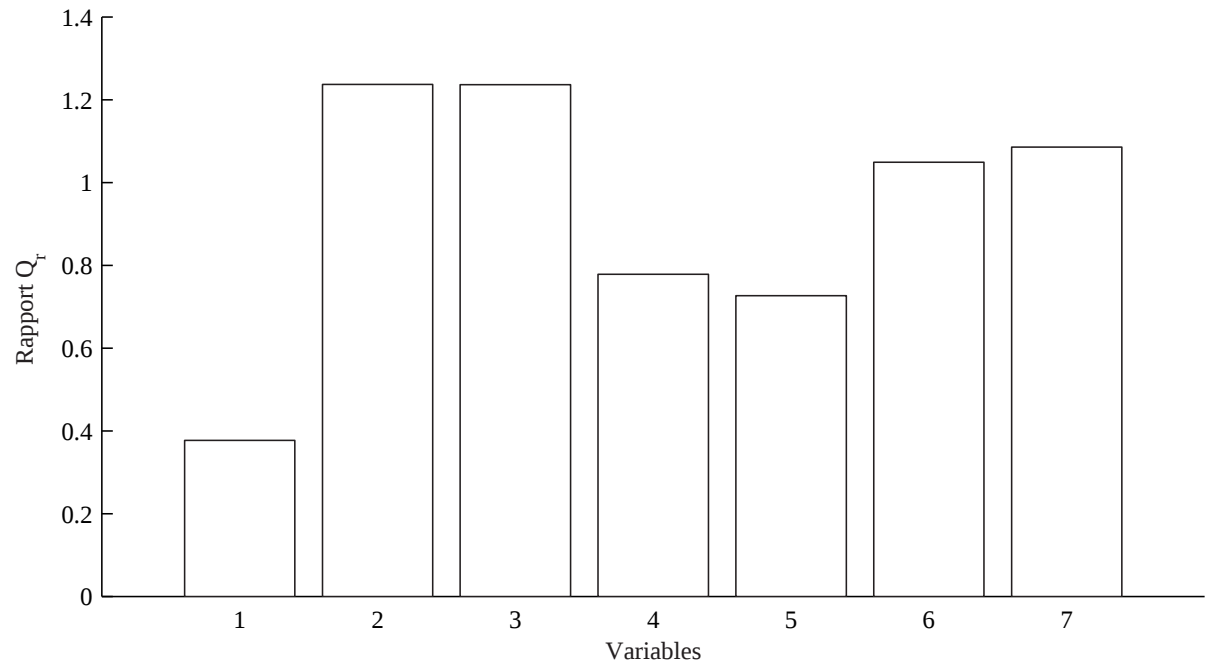
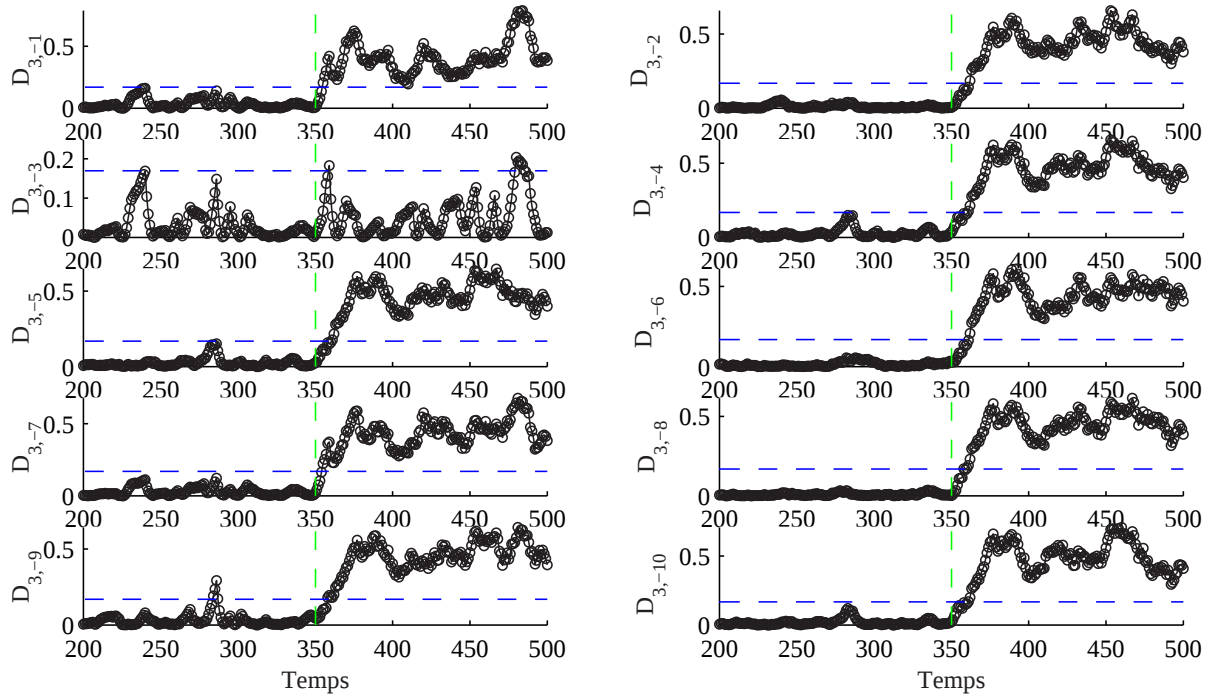


Figure 3.28 – Localisation de la variable  $x_1$  avec le rapport  $Q_r$ .

Figure 3.29 – Evolution des différents  $D_{i,-j}$  et localisation de la variable  $x_3$ 

$G_j$  permet de calculer le vecteur de mesures reconstruit  $\mathbf{x}_j$  en utilisant le vecteur d'entrée  $\mathbf{x}$  :

$$\mathbf{x}_j^T = \mathbf{x}^T G_j^T = [x_1 \dots z_j \dots x_m] \quad (3.126)$$

Ainsi, l'ensemble des indices de détection  $\{Res_i\}$  ( $i = 1, \dots, 5$ ) similaires au  $SPE$  et basé sur  $G_j$  sont définis :

$$Res_i = \|\mathbf{e}_i\|^2 = \|H\mathbf{x}\|^2 \quad (3.127)$$

où les différentes expressions de  $H$  sont représentées dans le tableau (tab 3.8).

| Résidus        | $H$                | Description                                                   |
|----------------|--------------------|---------------------------------------------------------------|
| $\mathbf{e}_1$ | $(I - \hat{C})$    | Mesures - Estimations par le modèle                           |
| $\mathbf{e}_2$ | $I - \hat{C}G_j$   | Mesures - Estimations par le modèle des mesures reconstruites |
| $\mathbf{e}_3$ | $G_j - \hat{C}$    | Mesures reconstruites - Estimations par le modèle             |
| $\mathbf{e}_4$ | $I - G_j$          | Mesures - Mesures reconstruites                               |
| $\mathbf{e}_5$ | $(I - \hat{C})G_j$ | Mesures reconstruites - Estimations des mesures reconstruites |

Table 3.8 – Description des différents expressions de  $H$  utilisées pour définir les résidus

A partir des différentes expression de  $H$ , nous pouvons écrire les équations des différents résidus en présence d'un défaut affectant la  $i^{eme}$  variable.

$$\begin{aligned}
\mathbf{e}_2 &= (I - CG_j) \mathbf{x} = (I - CG_j) (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & 1 & \dots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{pmatrix} - \begin{pmatrix} c_{11} & c_{12} & \dots & c_{1m} \\ \vdots & \vdots & \ddots & \vdots \\ c_{j1} & \vdots & c_{jj} & \dots & c_{jm} \\ \vdots & \vdots & \vdots & \ddots & \vdots \\ 0 & \dots & 0 & 1 \end{pmatrix} \end{pmatrix} (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} 1 - \left( c_{11} + \frac{c_{1j}^2}{1 - c_{jj}} \right) & \dots & -c_{j-1,1} - \frac{c_{1j}c_{j,j-1}}{1 - c_{jj}} & 0 & -c_{j+1,1} - \frac{c_{1j}c_{j,j+1}}{1 - c_{jj}} & \dots & -c_{1m} - \frac{c_{1j}c_{jm}}{1 - c_{jj}} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ -c_{j1} - \frac{c_{jj}c_{j1}}{1 - c_{jj}} & \dots & -c_{j-1,j} - \frac{c_{jj}c_{j,j-1}}{1 - c_{jj}} & 1 - c_{j+1,j} - \frac{c_{jj}c_{j,j+1}}{1 - c_{jj}} & \vdots & \vdots & -c_{jm} - \frac{c_{jj}c_{jm}}{1 - c_{jj}} \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ -c_{m1} - \frac{c_{mj}c_{j1}}{1 - c_{jj}} & \dots & \vdots & 0 & \vdots & \dots & 1 - \left( c_{mm} + \frac{c_{mj}^2}{1 - c_{jj}} \right) \end{pmatrix} (\mathbf{x}^* + \xi_i d)
\end{aligned}
\tag{3.129}$$

A partir de l'expression du vecteur des résidus  $\mathbf{e}_2$ , il est clair que le défaut se propage dans chaque résidu quelle que soit la variable reconstruite.

$$\begin{aligned}
\mathbf{e}_3 &= (G_j - C) \mathbf{x} = (G_j - C) (\mathbf{x}^* + \xi_i d) \\
&= \left( \begin{pmatrix} 1 & 0 & \cdot & \cdot & \cdot & \cdot & 0 \\ 0 & 1 & 0 & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \frac{c_{j1}}{1-c_{jj}} & \cdot & \frac{c_{j,j-1}}{1-c_{jj}} & 0 & \frac{c_{j,j+1}}{1-c_{jj}} & \cdot & \frac{c_{jm}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & 1 & 0 & \cdot \\ 0 & \cdot & \cdot & \cdot & \cdot & \cdot & 1 \end{pmatrix} - \begin{pmatrix} c_{11} & c_{12} & \cdot & \cdot & \cdot & c_{1m} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ c_{j1} & \cdot & c_{jj} & \cdot & c_{jm} & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ c_{m1} & \cdot & \cdot & \cdot & \cdot & c_{mm} \end{pmatrix} \right) (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} 1 - c_{11} & \cdot & -c_{1,j-1} & -c_{1j} & -c_{1,j+1} & \cdot & -c_{1m} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \frac{c_{j1}}{1-c_{jj}} - c_{j1} & \cdot & \frac{c_{j,j-1}}{1-c_{jj}} - c_{j-1,j} & -c_{jj} & \frac{c_{j,j+1}}{1-c_{jj}} - c_{j+1,j} & \frac{c_{jm}}{1-c_{jj}} - c_{jm} & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ -c_{m1} & \cdot & \cdot & -c_{mj} & \cdot & \cdot & 1 - c_{mm} \end{pmatrix} (\mathbf{x}^* + \xi_i d) \tag{3.130}
\end{aligned}$$

De même pour le vecteur des résidus  $\mathbf{e}_3$ , le défaut se propage dans tous les résidus même si la variable en défaut est la variable reconstruite ( $i = j$ ).

$$\begin{aligned}
\mathbf{e}_4 &= (I - G_j) \mathbf{x} = (I - G_j) (\mathbf{x}^* + \xi_i d) \\
&= \left( \begin{pmatrix} 1 & 0 & \cdot & \cdot & \cdot & \cdot & 0 \\ 0 & 1 & 0 & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \frac{c_{j1}}{1-c_{jj}} & \cdot & \frac{c_{j,j-1}}{1-c_{jj}} & 0 & \frac{c_{j,j+1}}{1-c_{jj}} & \cdot & \frac{c_{jm}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & 1 & 0 & \cdot \\ 0 & \cdot & \cdot & \cdot & \cdot & \cdot & 1 \end{pmatrix} - \begin{pmatrix} 1 & 0 & \cdot & \cdot & \cdot & \cdot & 0 \\ 0 & 1 & 0 & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \frac{c_{j1}}{1-c_{jj}} & \cdot & \frac{c_{j,j-1}}{1-c_{jj}} & 0 & \frac{c_{j,j+1}}{1-c_{jj}} & \cdot & \frac{c_{jm}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ 0 & \cdot & \cdot & \cdot & \cdot & \cdot & 1 \end{pmatrix} \right) (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} 0 & \cdot & 0 & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ -\frac{c_{j1}}{1-c_{jj}} & \cdot & -\frac{c_{j,j-1}}{1-c_{jj}} & 1 - \frac{c_{j,j+1}}{1-c_{jj}} & -\frac{c_{jm}}{1-c_{jj}} & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ 0 & \cdot & \cdot & 0 & \cdot & \cdot & 0 \end{pmatrix} (\mathbf{x}^* + \xi_i d) \tag{3.131}
\end{aligned}$$

$$\begin{aligned}
\mathbf{e}_5 &= (I - C) G_j \mathbf{x} = (I - C) G_j (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} 1 - c_{11} & -c_{12} & \cdot & \cdot & \cdot & -c_{1m} \\ -c_{21} & 1 - c_{22} & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ -c_{j1} & \cdot & \cdot & 1 - c_{jj} & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & -c_{m-1,m} \\ -c_{m1} & \cdot & \cdot & \cdot & -c_{m,m-1} & 1 - c_{mm} \end{pmatrix} \begin{pmatrix} 1 & 0 & \cdot & \cdot & \cdot & 0 \\ 0 & 1 & 0 & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & 1 & 0 & \cdot \\ \cdot & \cdot & \cdot & \cdot & 0 & \cdot \\ \frac{c_{j1}}{1-c_{jj}} & \cdot & \cdot & \frac{c_{j,j-1}}{1-c_{jj}} & 0 & \frac{c_{jm}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & 0 & 1 & \cdot \end{pmatrix} (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} (1 - c_{11}) - \frac{c_{1j}^2}{1-c_{jj}} & \cdot & \cdot & -c_{1,j-1} - \frac{c_{1j}c_{j,j-1}}{1-c_{jj}} & 0 & -c_{1,j+1} - \frac{c_{1j}c_{j,j+1}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ 0 & \cdot & \cdot & \cdot & 0 & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ -c_{m1} - \frac{c_{mj}c_{j1}}{1-c_{11}} & \cdot & \cdot & \cdot & 0 & \cdot \end{pmatrix} \begin{pmatrix} 0 & \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix} (\mathbf{x}^* + \xi_i d) \\
&= \begin{pmatrix} (1 - c_{11}) - \frac{c_{1j}^2}{1-c_{jj}} & \cdot & \cdot & -c_{1,j-1} - \frac{c_{1j}c_{j,j-1}}{1-c_{jj}} & 0 & -c_{1,j+1} - \frac{c_{1j}c_{j,j+1}}{1-c_{jj}} \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ 0 & \cdot & \cdot & \cdot & 0 & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ -c_{m1} - \frac{c_{mj}c_{j1}}{1-c_{11}} & \cdot & \cdot & \cdot & 0 & \cdot \end{pmatrix} \begin{pmatrix} 0 & \cdot & \cdot & \cdot & \cdot & 0 \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \\ \cdot & \cdot & \cdot & \cdot & \cdot & \cdot \end{pmatrix} (\mathbf{x}^* + \xi_i d)
\end{aligned} \tag{3.132}$$

L'expression du vecteur des résidus  $\mathbf{e}_5$  est très intéressante, car si la variable en défaut (la  $i^{eme}$ ) est la variable reconstruite ( $i = j$ ), le défaut est éliminé et aucun résidu n'est affecté.

Pour les différentes expressions des  $\mathbf{e}_i$  ( $i = 1, \dots, 5$ ), nous supposons que la reconstruction s'effectue dans la direction du défaut ( $i = j$ ). Nous pouvons constater que le défaut se propage dans les résidus  $\mathbf{e}_1$ ,  $\mathbf{e}_2$ ,  $\mathbf{e}_3$ ,  $\mathbf{e}_4$  et qu'il est éliminé de  $\mathbf{e}_5$  (3.132). La figure (fig 3.30) présente les différents indices calculés à partir de l'exemple 1 avec un défaut affectant la variable  $x_1$ .

La propriété de cet indice est très importante car elle est à la base d'une méthode de localisation de défaut proposée par Dunia et *al.* [14].

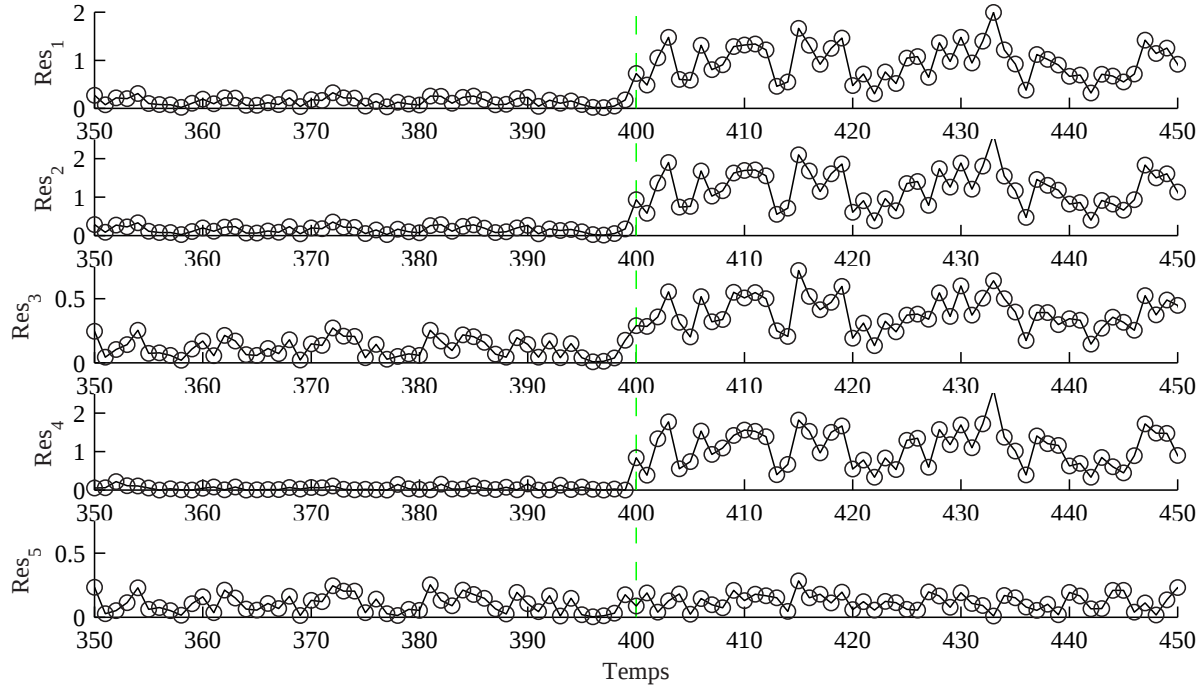


Figure 3.30 – Différents résidus obtenus à partir du modèle ACP de l'exemple 1.

Si on note par  $SPE_j$ , le  $SPE$  calculé après la reconstruction de la  $j^{eme}$  variable alors on peut écrire :

$$\begin{aligned} SPE_j &= \|\tilde{\mathbf{x}} - \tilde{d}_j \tilde{\xi}_j^\circ\|^2 \\ &= \|\tilde{\mathbf{x}}\|^2 - \tilde{d}_j^2 \\ &= SPE - \tilde{d}_j^2 \end{aligned} \quad (3.133)$$

Rappelons que  $\tilde{\mathbf{x}} = \mathbf{e}$ ,  $\tilde{d}_j$  est l'estimation du défaut dans la direction  $\xi_j$  projeté dans le sous-espace résiduel et que  $\tilde{\xi}_j^\circ = \tilde{\xi}/\|\tilde{\xi}\|$ .

$$SPE - SPE_j = \tilde{d}_j^2 \quad (3.134)$$

$\tilde{d}_j^2$  mesure la réduction dans le  $SPE$  due à la reconstruction dans la direction  $\xi_j$ . A partir de l'équation (3.134), on peut identifier le défaut par simple comparaison des amplitudes de  $SPE$  et  $SPE_j$ .

A partir des résidus calculés avant et après reconstruction, nous pouvons calculer un indice de validité du capteur (IVC) :

$$\eta_j^2 = \frac{SPE_j}{SPE} = 1 - \frac{\tilde{d}_j^2}{SPE} \quad (3.135)$$

où  $SPE$  est calculé avant reconstruction et  $SPE_j$  est calculé après la reconstruction de la valeur du  $j^{eme}$  capteur.

On peut rencontrer les trois situations suivantes :

1. le  $j^{eme}$  capteur est en défaut et le  $i^{eme}$  est reconstruit :

$$\eta_j^2 = 1 - \frac{(\tilde{\xi}_j^{\circ T} \tilde{\mathbf{x}})^2}{\|\tilde{\mathbf{x}}\|^2} = 1 - \frac{(\tilde{\xi}_j^{\circ T} (\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_i^{\circ}))^2}{\|\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_i^{\circ}\|^2} = 1 - \frac{(\tilde{\xi}_j^{\circ T} \tilde{\mathbf{x}}^*)^2}{\|\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_i^{\circ}\|^2} \quad (3.136)$$

l'indice  $\eta_j$  tend vers 1.

2. le capteur en défaut est le capteur reconstruit :

$$\eta_j^2 = 1 - \frac{(\tilde{\xi}_j^{\circ T} \tilde{\mathbf{x}})^2}{\|\tilde{\mathbf{x}}\|^2} = 1 - \frac{(\tilde{\xi}_j^{\circ T} (\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_j^{\circ}))^2}{\|\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_j^{\circ}\|^2} = 1 - \frac{(\tilde{\xi}_j^{\circ T} \tilde{\mathbf{x}}^* + \tilde{d})^2}{\|\tilde{\mathbf{x}}^* + \tilde{d}\tilde{\xi}_j^{\circ}\|^2} \quad (3.137)$$

et l'indice  $\eta_j$  tend vers zéro.

3. pas de capteur défaillant, dans ce cas nous avons :

$$\eta_j^2 = 1 - \frac{(\tilde{\xi}_j^{\circ T} \tilde{\mathbf{x}}^*)^2}{\|\tilde{\mathbf{x}}^*\|^2} = \frac{\tilde{\mathbf{x}}^{*T} (I - \tilde{\xi}_j^{\circ} \tilde{\xi}_j^{\circ T}) \tilde{\mathbf{x}}^*}{\tilde{\mathbf{x}}^{*T} \tilde{\mathbf{x}}^*} \quad (3.138)$$

l'indice  $\eta_j$  peut prendre des valeurs entre zéro et un, et en particulier  $\eta_j = 0$  si  $\tilde{\mathbf{x}}^* = \|\tilde{\mathbf{x}}^*\| \tilde{\xi}_j^{\circ}$ . Donc, l'indice  $\eta_j$  ne peut pas être utilisé pour la détection mais uniquement pour la localisation.

La figure (fig 3.31) présente l'évolution de l'indice de localisation pour les différentes variables dans le cas de l'exemple 1, avec un défaut affectant la variable  $x_1$  à partir de l'instant 400. L'indice correspondant à la variable  $x_1$  est proche de zéro ce qui indique que c'est la variable incriminée.

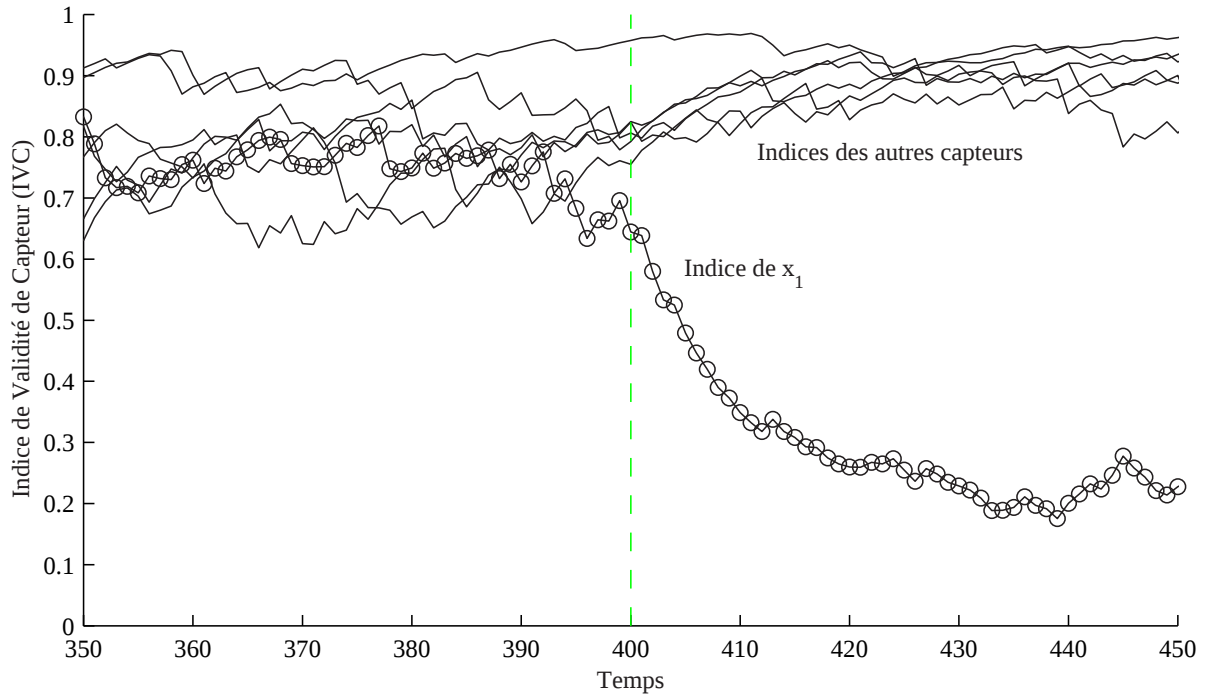


Figure 3.31 – Indices de validité de capteur filtrés avec un défaut affectant  $x_1$  entre les instants 400 et 500

### 3.3.2.4 Conditions nécessaires et suffisantes pour la localisation par reconstruction

Nous avons déjà montré que la reconstruction annule l'erreur d'estimation de la variable reconstruite. Ainsi, si on suppose que l'espace résiduel est de dimension  $(m - \ell) = 1$ , alors l'expression du vecteur des résidus est donnée par :

$$\mathbf{e} = p_m t_m \quad (3.139)$$

où  $t_m$  est la dernière composante principale et  $p_m$  est le dernier vecteur propre. Rappelons que  $p_m = [p_{1m} p_{2m} \dots p_{mm}]^T$ , alors tous les résidus sont identiques à un facteur  $p_{im}$  près. Ainsi, si on reconstruit dans n'importe quelle direction le résidu correspondant est nul et  $\mathbf{e}_j$  représentant l'erreur de reconstruction de la variable  $j$  est nul  $\forall j$ . D'où la condition nécessaire de localisation par reconstruction donnée par :

$$(m - \ell) \geq 2 \quad (3.140)$$

Pour illustrer cette condition nous allons supposer que la dimension de l'espace résiduel est  $(m - \ell) = 1$  et que nous avons le dernier vecteur propre suivant :  $\tilde{p}^T = \frac{1}{\sqrt{a^2 + b^2 + c^2}} [a \ b \ c]$ . Les différentes directions de défauts dans le cas d'un processus tridimensionnel sont données par :

$$[\xi_1 \ \xi_2 \ \xi_3] = \begin{bmatrix} 1 & 0 & 0 \\ 0 & 1 & 0 \\ 0 & 0 & 1 \end{bmatrix}$$

$$\tilde{C} = \tilde{p}\tilde{p}^T = \frac{1}{(a^2 + b^2 + c^2)} \begin{bmatrix} a^2 & ab & ac \\ ab & b^2 & bc \\ ac & bc & c^2 \end{bmatrix} = \frac{1}{\sqrt{a^2 + b^2 + c^2}} [a\tilde{p} \ b\tilde{p} \ c\tilde{p}]$$

$$[\tilde{\xi}_1 \ \tilde{\xi}_2 \ \tilde{\xi}_3] = \tilde{C} [\xi_1 \ \xi_2 \ \xi_3] = \frac{1}{\sqrt{a^2 + b^2 + c^2}} [a\tilde{p} \ b\tilde{p} \ c\tilde{p}]$$

Après normalisation, nous obtenons l'équation suivante :

$$[\tilde{\xi}_1^\circ \ \tilde{\xi}_2^\circ \ \tilde{\xi}_3^\circ] = [\tilde{p} \ \tilde{p} \ \tilde{p}]$$

Ainsi, le  $SPE$  après reconstruction peut s'écrire sous la forme :

$$SPE_i = \|\mathbf{x}_i - \hat{\mathbf{x}}_i\|^2 = \|(I - \hat{C}) \mathbf{x}_i\|^2 \quad (3.141)$$

A partir de l'équation (3.141), l'expression du  $SPE$  après reconstruction dans la  $i^{eme}$  direction est donnée par :

$$SPE_i = \|(I - \tilde{\xi}_i^\circ \tilde{\xi}_i^{\circ T}) \tilde{\mathbf{x}}\|^2$$

Puisque  $\tilde{\mathbf{x}} \in S_r$ , on peut écrire  $\tilde{\mathbf{x}} = \|\tilde{\mathbf{x}}\| \tilde{p}$  et l'équation précédente prend la forme suivante :

$$SPE_i = \|(I - \tilde{p}\tilde{p}^T) \tilde{p}\|^2 \|\tilde{\mathbf{x}}\|^2 = 0$$

car  $(I - \tilde{p}\tilde{p}^T)\tilde{p} = 0$

Ainsi,  $SPE_i = 0$  pour toutes les directions  $\forall \|\tilde{\mathbf{x}}\|$ .

De plus, si les défauts dans deux directions  $i$  et  $j$  sont tels que  $\tilde{\xi}_i^\circ = \pm\tilde{\xi}_j^\circ$ , alors on aura :

$$SPE_j = \|(I - \tilde{\xi}_j^\circ \tilde{\xi}_j^{\circ T})\tilde{\mathbf{x}}\|^2 = \|(I - \tilde{\xi}_i^\circ \tilde{\xi}_i^{\circ T})\tilde{\mathbf{x}}\|^2 = SPE_i$$

et les défauts  $d_i$  et  $d_j$  ne sont pas localisables.

D'où la condition suffisante de localisation par reconstruction donnée par :

$$\xi_i^\circ \neq \pm \xi_j^\circ \quad \forall i \neq j \quad (3.142)$$

Cette condition indique que les directions des défauts projetées dans le sous-espace des résidus ne doivent pas être colinéaires.

### 3.3.2.5 Localisation par reconstruction utilisant l'indice $D_i$

Nous avons montré précédemment, sur un exemple simple, les limites de l'utilisation de l'indice  $SPE$  pour la détection de défauts. Ainsi, la procédure de localisation basée sur le principe de reconstruction utilisant le  $SPE$  n'est plus exploitable. Pour cette raison, une approche de localisation basée sur l'indice de détection  $D_i$  filtré et le principe de reconstruction est proposée. L'analyse de l'indice de détection  $D_i$  filtré, avant et après reconstruction, permet de localiser la variable incriminée. Pour plus d'explications, nous allons développer l'expression de l'indice  $D_i$  calculé après la reconstruction de la  $j^{eme}$  variable et qui sera noté par  $D_i^{(j)}$ .

Pour étudier la propagation d'un défaut sur l'indice  $D_i^{(j)}$ , nous allons dans un premier temps développer l'expression du vecteur de mesure  $\mathbf{x}_j$  dont la  $j^{eme}$  variable est reconstruite (2.41).

$$\mathbf{x}_j(k) = \begin{bmatrix} \mathbf{x}_{-j}(k) \\ z_j(k) \\ \mathbf{x}_{+j}(k) \end{bmatrix} = \begin{bmatrix} \mathbf{x}_{-j}(k) \\ 0 \\ \mathbf{x}_{+j}(k) \end{bmatrix} + \xi_j z_j(k) \quad (3.143)$$

et

$$\begin{bmatrix} \mathbf{x}_{-j}(k) \\ 0 \\ \mathbf{x}_{+j}(k) \end{bmatrix} = \text{diag}(\zeta_j) \mathbf{x}(k) \quad (3.144)$$

où  $\mathbf{x}(k) = \mathbf{x}^*(k) + \xi_f d(k)$  est le vecteur de mesure avec un défaut  $d(k)$  dans la direction  $\xi_f$ .  $\xi_j$  est la direction de reconstruction et  $\text{diag}(\zeta_j)$  est une matrice diagonale contenant un 0 dans la  $j^{eme}$  position et des 1 ailleurs. Ainsi, l'expression de  $D_i^{(j)}(k)$  est donnée par :

$$D_i^{(j)}(k) = \sum_{l=m-i+1}^m t_l^2(k) = \sum_{l=m-i+1}^m \left\{ p_l^T \left( \text{diag}(\zeta_j) + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \right) (\mathbf{x}^*(k) + \xi_f d(k)) \right\}^2 \quad (3.145)$$

Le développement de l'expression à l'intérieur de la somme permet d'écrire :

$$p_l^T \left( \text{diag}(\zeta_j) + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \right) (\mathbf{x}^*(k) + \xi_f d(k)) =$$

$$p_l^T \left( \text{diag}(\zeta_j) + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \right) \mathbf{x}^*(k) \quad (3.146)$$

$$+ p_l^T \left( \text{diag}(\zeta_j) \xi_f + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \xi_f \right) d(k)$$

Si on considère le cas où la variable en défaut est la variable reconstruite ( $j = f$ ), alors :

$$\left\{ \begin{array}{l} \text{diag}(\zeta_j) \xi_j = \begin{pmatrix} 1 & 0 & \dots & 0 \\ 0 & . & & . \\ . & 1 & & . \\ . & 0 & & . \\ . & . & 1 & . \\ . & . & . & 0 \\ 0 & . & . & 0 & 1 \end{pmatrix} \begin{pmatrix} 0 \\ . \\ 0 \\ 1 \\ 0 \\ . \\ 0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ . \\ . \\ . \\ . \\ 0 \end{pmatrix} \\ \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \xi_j = \begin{pmatrix} 0 \\ 0 \\ . \\ . \\ . \\ . \\ 0 \end{pmatrix} \end{array} \right. \quad (3.147)$$

Ainsi, l'expression finale de l'indice  $D_i^{(j)}$  est donnée par :

$$D_i^{(j)}(k) = \left\{ \begin{array}{l} \sum_{l=m-i-1}^m \left\{ p_l^T \left( \text{diag}(\zeta_j) + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \right) \mathbf{x}^*(k) \right\}^2 \text{ pour } j = f \\ \sum_{l=m-i+1}^m \left\{ p_l^T \left( \text{diag}(\zeta_j) + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \right) \mathbf{x}^*(k) \right. \\ \left. + p_l^T \left( \text{diag}(\zeta_j) \xi_f + \xi_j \frac{\begin{bmatrix} c_{-j}^T & 0 & c_{+j}^T \end{bmatrix}}{1 - c_{jj}} \xi_f \right) d(k) \right\}^2 \text{ pour } j \neq f \end{array} \right. \quad (3.148)$$

A partir de l'équation (3.148), on constate que si la reconstruction se fait dans la direction du défaut ( $j = f$ ), l'effet du défaut sera éliminé et ainsi on pourra localiser la variable incriminée. L'indice filtré  $D_i^{(j)}$  qui se trouve en dessous du seuil de détection indique que la  $j^{eme}$  variable est en défaut. On peut également utiliser un filtre EWMA pour filtrer l'indice  $D_i^{(j)}$  comme dans le cas du *SPE*. On notera l'indice filtré par  $\bar{D}_i^{(j)}$ .

Ainsi, nous pouvons définir l'indice de localisation suivant à l'instant  $k$  :

$$A_i^{(j)}(k) = \frac{\bar{D}_i^{(j)}(k)}{\bar{\tau}_{i,\alpha}^2} \quad (3.149)$$

où  $\bar{\tau}_{i,\alpha}^2$  est le seuil de détection du  $\bar{D}_i$ .

Si la variable  $j$  est la variable incriminée alors l'indice  $A_i^{(j)}$  sera inférieur à 1 :

$$\begin{cases} A_i^{(j)} > 1 \text{ pour } j \neq f \\ A_i^{(j)} \leq 1 \text{ pour } j = f \end{cases} \quad (3.150)$$

Pour l'exemple 2, en appliquant la procédure de localisation par reconstruction à l'indice de détection  $\bar{D}_3$  sur lequel on a détecté le défaut affectant la variable  $x_3$ , les résultats sont représentés sur la figure (fig 3.32). L'indice  $A_3^{(3)}$  (indice calculé après reconstruction de la variable  $x_3$ ) est en dessous du seuil de détection ce qui indique que la variable  $x_3$  est la variable incriminée (fig 3.33).

De même pour le défaut affectant la variable  $x_9$  et qui a été détecté avec l'indice  $\bar{D}_1$ , la figure (fig 3.34) présente l'évolution de l'indice  $D_1$  après reconstruction des différents variables et la figure (fig 3.35) présente l'indice  $A_1^{(j)}$  calculé à partir de l'indice  $\bar{D}_1^{(j)}$ . L'indice  $A_1^{(9)}$  est inférieur à 1, ce qui indique que la variable  $x_9$  est la variable incriminée.

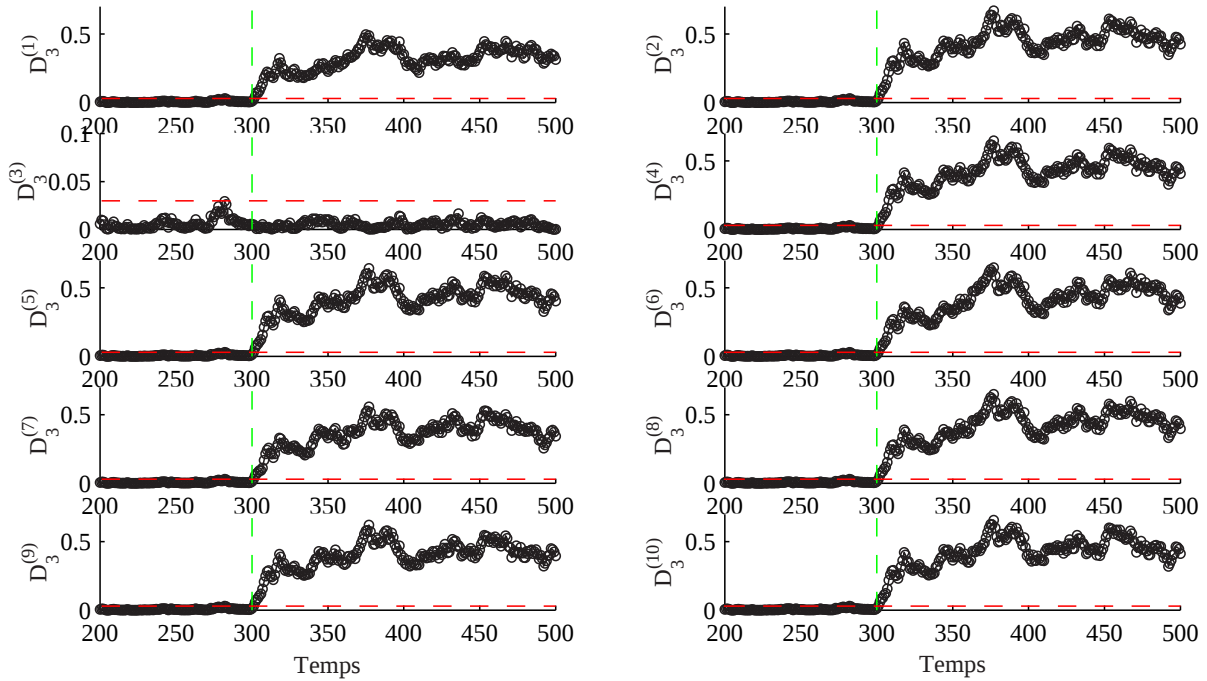


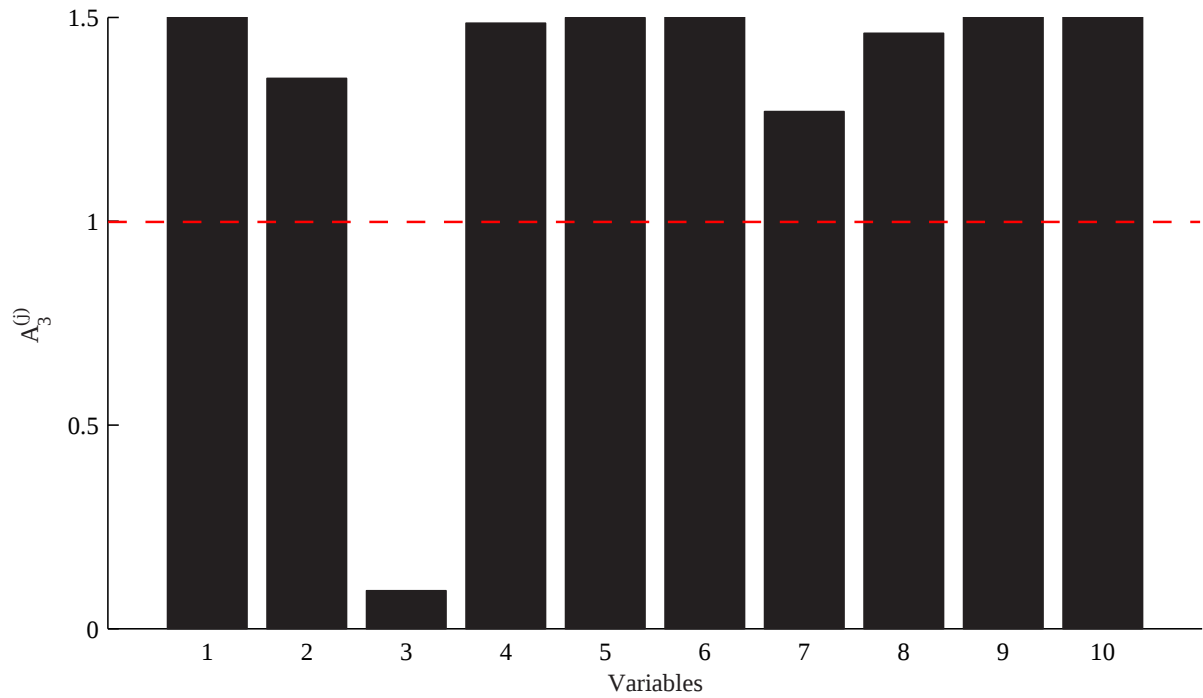
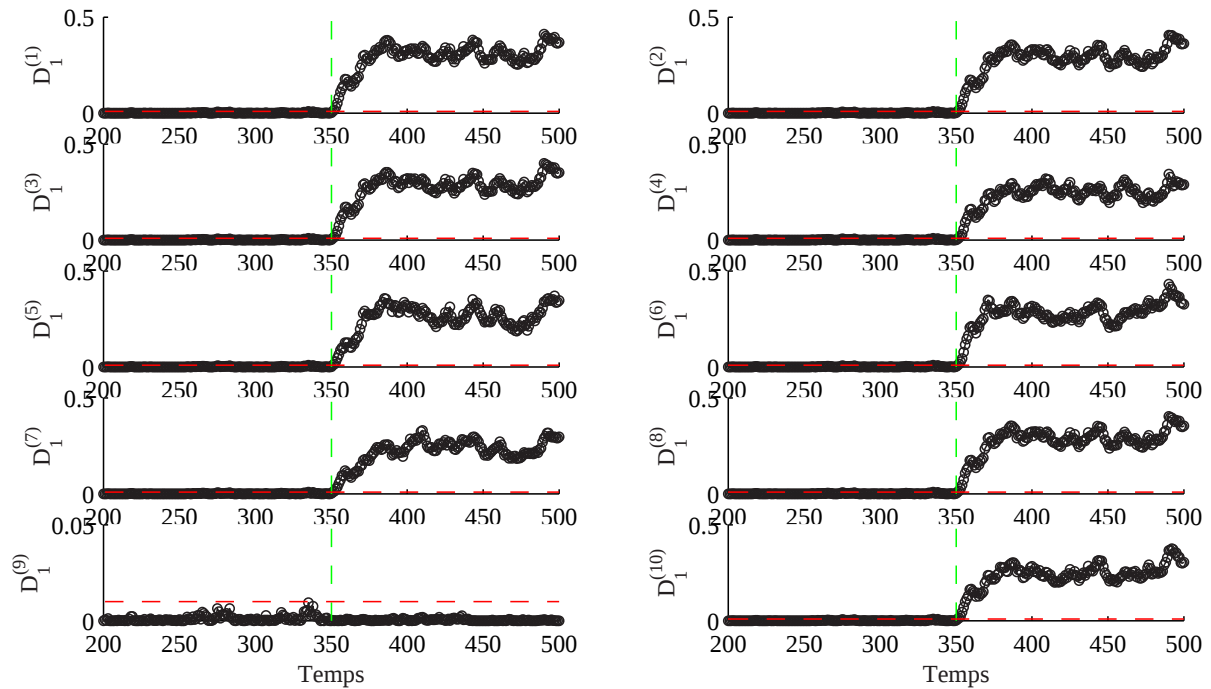
Figure 3.32 – Evolution des différents indices  $D_3$  filtrés après reconstruction des différents variables (exemple 2)

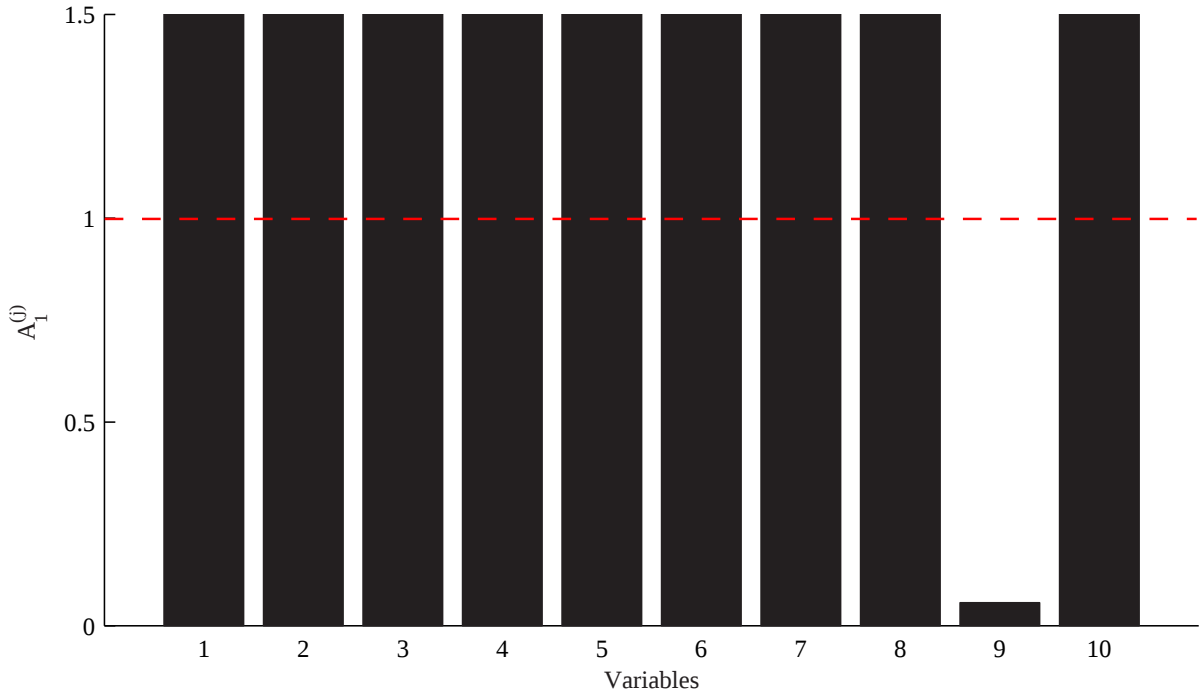
### 3.3.3 Localisation par calcul des contributions

Le calcul des contributions de chaque variable à la valeur de l'indice de détection est une approche utilisée pour la localisation des défauts, la variable ayant la plus forte contribution est considérée comme la variable en défaut. Cette approche est largement utilisée dans le cadre de l'ACP [113, 55, 68, 93, 108, 50, 7].

Dans cette section nous allons présenter deux définitions des contribution à l'indice proposé  $D_i$ . Ces deux définitions sont basées sur les définitions des contribution à l'indice  $SPE$  et à la statistique  $T^2$  [113, 55].

Dans le cas de l'indice de détection  $SPE$ , la contribution  $cont_j^{SPE}(k)$  de la  $j^{eme}$  variable à l'instant  $k$  est définie par l'équation [113, 55] :

Figure 3.33 – Indices  $A_3^{(j)}$  et localisation de la variable en défaut  $x_3$ Figure 3.34 – Evolution des différents indices  $D_1$  filtrés après reconstruction des différents variables (exemple 2)

Figure 3.35 – Indices  $A_1^{(j)}$  et localisation de la variable en défaut  $x_9$ 

$$cont_j^{SPE}(k) = (e_j(k))^2 = (x_j(k) - \hat{x}_j(k))^2 \quad (3.151)$$

où  $x_j(k)$  est le  $j^{eme}$  élément du vecteur de mesures  $\mathbf{x}$ .

Si on considère le vecteur de mesures suivant dont la  $i^{eme}$  composante est affectée par un défaut  $d$  :

$$\mathbf{x} = \mathbf{x}^* + \xi_i d \quad (3.152)$$

Le vecteur des résidus peut être obtenu par l'équation suivante :

$$\mathbf{e}(k) = (I - \hat{P}\hat{P}^T) \mathbf{x}(k) = \tilde{C} (\mathbf{x}^*(k) + \xi_i d(k)) \quad (3.153)$$

Le  $SPE$  peut être calculé par :

$$SPE(k) = \mathbf{e}^T(k) \mathbf{e}(k) = \sum_{j=1}^m e_j^2(k) \quad (3.154)$$

A partir de l'équation (3.153), la quantité  $e_j$  peut être exprimée par :

$$e_j(k) = \sum_{l=1}^m \tilde{c}_{jl} x_l^*(k) + \tilde{c}_{ji} d(k) \quad (3.155)$$

Ainsi, la contribution de la  $j^{eme}$  variable au  $SPE$  sera définie comme :

$$cont_j^{SPE}(k) = \left( \sum_{l=1}^m (\tilde{c}_{jl} x_l^*(k) + \tilde{c}_{ji} d(k)) \right)^2 \quad (3.156)$$

Ainsi, les contributions de toutes les variables sont affectées par le défaut.

Cette dernière expression peut être utilisée pour le calcul des contributions à l'indice proposé  $\bar{D}_i$  (3.81). En tenant compte du fait que :

$$\bar{D}_i(k) = \sum_{j=m-i+1}^m \bar{t}_j^2(k) = \overline{SPE}_{(i)}(k) \quad (3.157)$$

où  $\overline{SPE}_{(i)}(k)$  est l'erreur quadratique filtrée calculée à partir d'un modèle ACP à  $(m-i)$  composantes, ce qui nous permet d'écrire :

$$\bar{D}_i(k) = \bar{\mathbf{e}}_{(i)}^T \bar{\mathbf{e}}_{(i)} = \sum_{j=1}^m \bar{e}_{j,(i)}^2(k) \quad (3.158)$$

où  $\mathbf{e}_{(i)}$  est le vecteur des résidus obtenu à partir d'un modèle ACP avec  $\ell = (m-i)$  composantes et  $e_{j,(i)}(k)$  est le  $j^{eme}$  élément de  $\mathbf{e}_{(i)}(k)$ . Ainsi, nous définissons les contributions des variables au  $\bar{D}_i(k)$  par :

$$cont_j^{\bar{D}_i}(k) = e_{j,(i)}^2(k) \quad (3.159)$$

La variable ayant la plus forte contribution par rapport aux autres variables est considérée comme la variable en défaut.

Le résultat de l'application de l'analyse des contributions aux indices de détection  $\bar{D}_3$  et  $\bar{D}_1$ , sur lesquels les défauts affectant, respectivement, les variables  $x_3$  et  $x_9$  (exemple 2) ont été détectés, est illustré sur les deux figures (fig 3.36) et (fig 3.37)

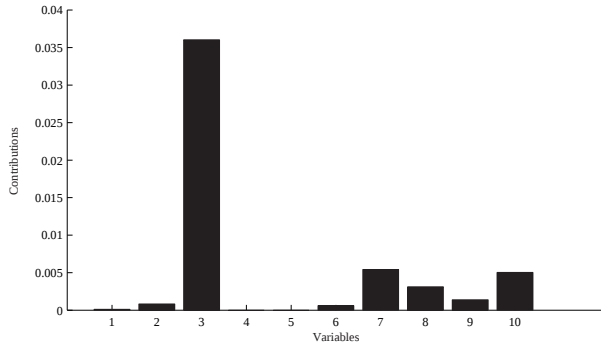


Figure 3.36 – Contributions des variables à l'indice  $\bar{D}_3$  avec un défaut affectant  $x_3$

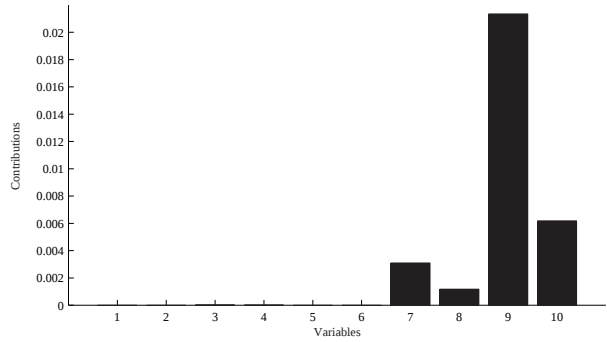


Figure 3.37 – Contributions des variables à l'indice  $\bar{D}_1$  avec un défaut affectant  $x_9$

Comme l'expression de l'indice  $D_i$  peut être exprimée en fonction des dernières composantes principales (3.76), la deuxième définition que nous proposons est basée sur le calcul des contributions des variables aux composantes principales. Dans un premier temps nous allons présenter les définitions des contribution à la statistique  $T^2$  car ces définition utilisent les contribution aux composantes principales.

Chaque composante s'exprime comme suit :

$$t_i(k) = p_i^T \mathbf{x}(k) = \sum_{j=1}^m p_{ij} x_j(k) \quad (3.160)$$

où  $p_i$  désigne le vecteur propre correspondant à la valeur propre  $\lambda_i$ .

Comme la statistique  $T^2$  n'est pas utilisée comme un indice de détection, le calcul des contributions à cette statistique n'a pas d'intérêt particulier. Cependant, nous allons le présenter à titre indicatif et l'exploiter pour calculer les contributions des variables à l'indice de détection  $D_i$  que nous avons proposé. Plusieurs auteurs se sont intéressés au calcul des contributions à la statistique  $T^2$  [113, 55, 68, 108].

Wise et Gallagher [113] proposent d'évaluer les contributions des variables à la statistique  $T^2$ . En utilisant l'équation (3.160), la contribution globale des variables à la  $i^{ieme}$  composante principale normalisée  $\left(\frac{t_i}{\sigma_i}\right)^2$  est définie par :

$$cont_i(k) = \frac{t_i(k)}{\sigma_i^2} \sum_{j=1}^m p_{ij} x_j(k) \quad (3.161)$$

où  $p_{ij}$  est la  $j^{eme}$  composante du vecteur propre  $p_i$  et  $\sigma_i = \sqrt{\lambda_i}$ .

On approxime la contribution d'une variable  $x_j$  à cette composante principale normalisée  $\left(\frac{t_i(k)}{\sigma_i}\right)^2$  par :

$$cont_{ij}(k) = \frac{t_i(k)}{\lambda_i} p_{ij} x_j(k) \quad (3.162)$$

La contribution totale de la variable  $x_j(k)$  à la statistique  $T^2$ , à l'instant  $k$ , sera donnée par l'équation :

$$Cont_j(k) = \sum_{i=1}^{\ell} cont_{ij}(k) \quad (3.163)$$

A partir de cette formule de calcul des contributions, on peut remarquer que chaque contribution dépend des termes croisés de la variable considérée avec les autres variables. Ainsi, cette définition ne représente pas vraiment la contribution de la variable à la statistique  $T^2$  mais une simple approximation, vue que la contribution exacte ne peut pas être calculée.

Kourti et al. [55] proposent une utilisation simultanée des composantes principales et des contributions des variables initiales. Lors de la détection d'un défaut, ils proposent d'analyser les composantes principales normalisées ayant subi une variation significative. La contribution totale de la variable  $x_j(k)$  sur les  $q$  composantes les plus élevées (parmi les  $\ell$  premières) est donnée par :

$$Cont_j(k) = \sum_{i=1}^q cont_{ij}(k) \quad (3.164)$$

avec les conditions suivantes sur les contributions :

- une contribution est mise à zéro si son signe est négatif (son signe est opposé au signe de la composante  $t_i$ ) ; on ne considère que les variables corrélées positivement à la composante (contribution à l'accroissement de la composante),
- les variables ayant la plus forte contribution sont celles qui sont en défaut.

Il faut noter que cette procédure de localisation est appliquée aux  $\ell$  premières composantes principales.

Par simple développement de l'expression de la  $i^{eme}$  composante principale, on peut exprimer la contribution de la  $j^{eme}$  variable à cette composante par l'équation [108] :

$$cont_{ij}(k) = \sum_{r=1}^m p_{ij} x_j(k) p_{ir} x_r(k) \quad (3.165)$$

| $p_4$  | $p_5$  | $p_6$  | $p_7$  | $p_8$  | $p_9$  | $p_{10}$ |
|--------|--------|--------|--------|--------|--------|----------|
| -0.017 | 0.445  | -0.189 | 0.737  | -0.045 | 0.023  | -0.002   |
| -0.039 | 0.564  | -0.077 | -0.664 | -0.126 | 0.002  | 0.003    |
| -0.059 | -0.267 | 0.068  | -0.060 | 0.807  | -0.216 | -0.021   |
| 0.030  | 0.162  | 0.687  | 0.061  | 0.002  | 0.015  | 0.013    |
| 0.079  | -0.167 | -0.679 | -0.077 | 0.006  | -0.007 | -0.007   |
| -0.773 | 0.171  | -0.061 | 0.017  | 0.102  | -0.044 | -0.005   |
| -0.191 | -0.417 | 0.075  | -0.023 | -0.212 | 0.675  | 0.311    |
| 0.554  | 0.249  | -0.038 | -0.010 | 0.303  | 0.320  | 0.192    |
| 0.145  | -0.207 | 0.067  | -0.004 | -0.226 | 0.001  | -0.819   |
| 0.152  | -0.218 | 0.059  | 0.016  | -0.362 | -0.626 | 0.440    |

Table 3.9 – Matrice des derniers vecteurs propres de la matrice de corrélation de l'exemple 2

Cependant, on remarque bien que les différentes expressions de  $cont_{ij}(k)$  (contribution de la variable  $j$  à la composante  $i$ ) dépendent directement de l'amplitude de la variable  $x_j$  et du coefficient  $p_{ij}$ . De ce fait, on ne peut mettre des seuils sur les contributions, que ce soit dans le cas du  $T^2$  ou du  $SPE$ , car ces contributions dépendent directement des amplitudes des variables. Ainsi, la localisation des variables en défaut, par analyse des contributions consiste à considérer les variables ayant la plus grande contribution à l'indice de détection comme des variables en défaut.

Les définitions des contributions données précédemment sont appliquées aux premières composantes principales. Pour les exploiter dans le cas de notre indice, il faut les appliquer aux dernières composantes. Ainsi, pour tester l'application de ces définition à notre indice de détection  $D_i$ , nous allons simuler un défaut (exemple 2) de telle sorte que la variable en défaut soit proche de zéro, ce qui peut toujours arriver (panne totale d'un capteur). Le défaut est simulé sur la variable  $x_1$  de l'exemple 2 entre les instants 430 et 500 avec une amplitude d'environ 20% de la plage de variation de cette variable.

La figure (fig 3.38) présente l'évolution des variables  $x_1$ ,  $x_2$  et le défaut affectant  $x_1$ , le défaut a été détecté sur l'indice  $D_4$  à l'instant 434 (fig 3.39). Le calcul des contributions à  $D_4$  à cet instant est présenté sur les figures (fig 3.40) et (fig 3.41). Prenons par exemple le cas de l'approche de Kourti et *al.* [55], la variable qui contribue le plus à cet indice est la variable  $x_2$  alors que c'est  $x_1$  qui est en défaut (fig 3.40). Pour préciser ce point, analysons l'expression de l'indice  $D_4$  :

$$D_4(k) = t_{10}^2(k) + t_9^2(k) + t_8^2(k) + t_7^2(k) \quad (3.166)$$

A partir de cette expression et en analysant la matrice des derniers vecteurs propres (tab 3.9), il est clair que tous les coefficients avec lesquels  $x_1$  intervient dans le calcul des composantes  $t_{10}$ ,  $t_9$  et  $t_8$  sont très faibles. Ainsi, le défaut est porté par la composante  $t_7$  car le vecteur  $p_7$  (septième colonne de la matrice des vecteurs propres) (tab 3.9) est à la base de la détection avec l'indice  $D_4$  puisque les coefficients correspondant aux variables  $x_1$  et  $x_2$  sont les plus significatifs sur ce vecteur.

On remarque bien que les coefficients des deux variables  $x_1$  et  $x_2$  sont presque égaux et de signes opposés puisque les deux variables sont corrélées, ce qui explique en fait qu'en absence de défaut les projections des deux variables sur le vecteur  $p_7$  se compensent. La présence d'un défaut affectant l'une des variables est facilement détectable par projection sur ce vecteur (incohérence de la relation entre les deux variables). Cependant, le calcul des contributions dépend directement

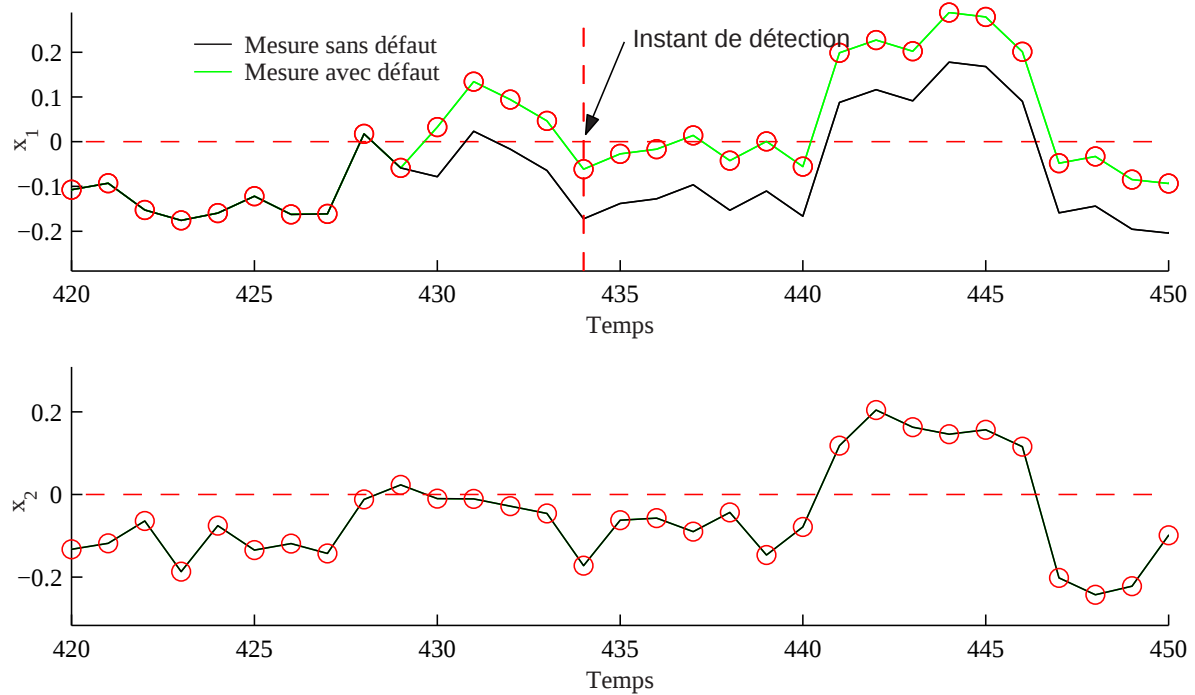


Figure 3.38 – Evolution des deux variables  $x_1$  et  $x_2$  avec un défaut affectant  $x_1$  à l'instant 430 et détecté à l'instant 434.

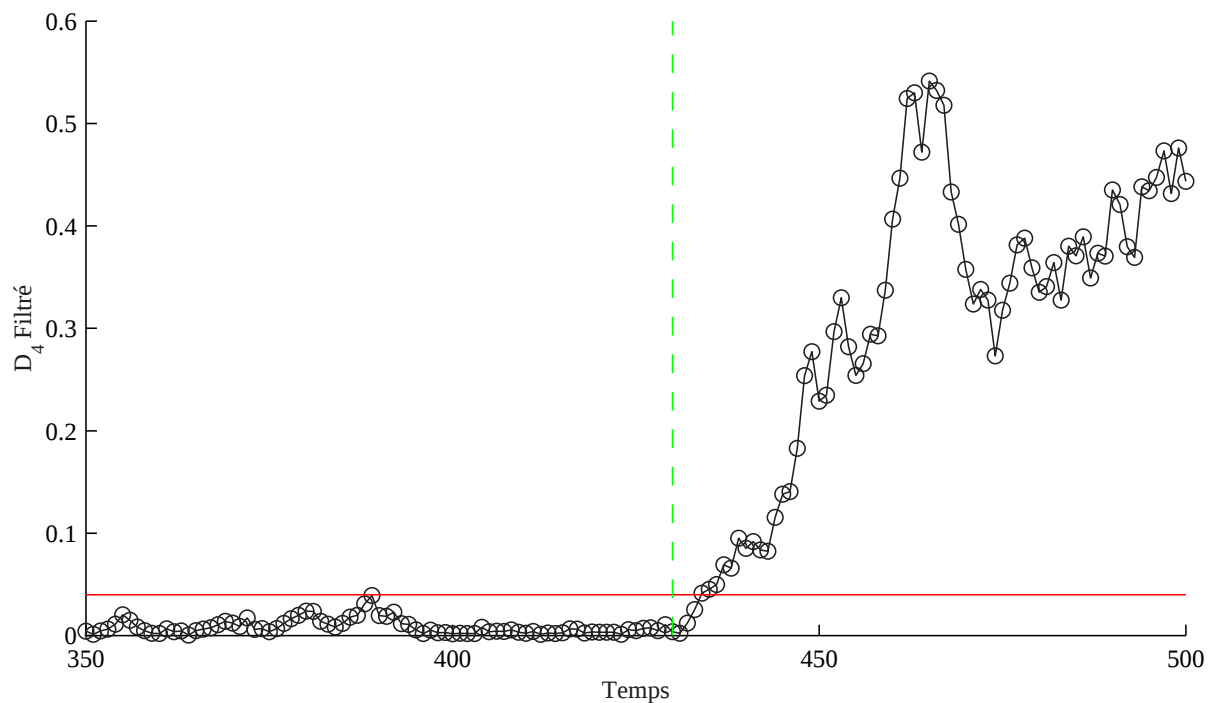


Figure 3.39 – Evolution de l'indice  $D_4$  filtré avec un défaut affectant  $x_1$  à l'instant 430.

des amplitudes des variables. Comme dans notre cas l'amplitude de la variable  $x_1$  affectée par le défaut est plus petite que celle de  $x_2$ , cela signifie que la contribution de la variable  $x_2$  sera plus

importante que celle de  $x_1$ . Avec cet exemple simple, nous avons mis en évidence les limites de la méthode de localisation par calcul des contributions aux dernières composantes principales en utilisant les approches existantes.

Pour résoudre ce problème, nous proposons d'analyser les dernières composantes principales (intervenant dans le calcul de l'indice  $D_i$ ) ayant subi des variations significatives, et de calculer la variation des contributions des différentes variables à ces composantes avant et après détection du défaut. Ainsi, en supposant qu'il n'y a pas de changement de régime de fonctionnement, les variations des contributions à la  $i^{eme}$  composante est définie par :

$$vcont_{ij} = \frac{1}{k_w} \left( \sum_{k=k_d-k_w}^{k_d-1} x_j(k) - \sum_{k=k_d}^{k_d+k_w-1} x_j(k) \right) p_{ij} \quad (3.167)$$

où  $k_d$  représente l'instant de détection du défaut et  $k_w$  représente un nombre de points sur lequel on doit calculer la moyenne de la variable avant et après détection.

Dans le cas où nous avons plusieurs composantes, soit  $q$ , qui présentent des variations significatives dues au défaut, les variations des contributions seront calculées en sommant les carrés des variations de chaque composante :

$$Vcont_j = \sum_{i=1}^q vcont_{ij}^2 \quad (3.168)$$

L'application de l'approche proposée à l'indice  $D_4$ , avec  $k_w = 5$ , permet de déterminer la variable incriminée comme le montre la figure (fig 3.42).

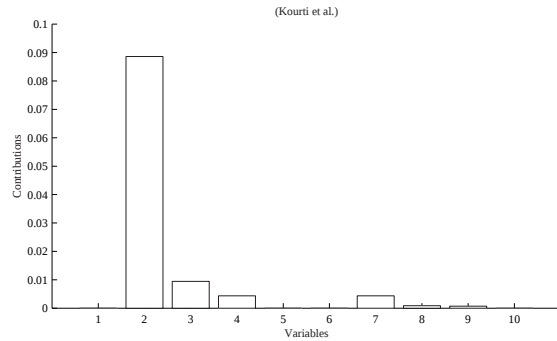


Figure 3.40 – Contributions des différentes variables à l'indice  $D_4$  à l'instant 434, approche de Kourti

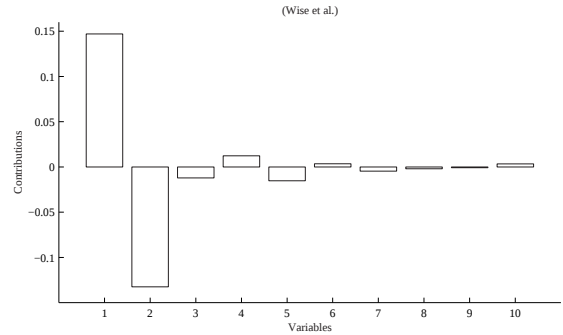


Figure 3.41 – Contributions des différentes variables à l'indice  $D_4$  à l'instant 434, approche de Wise

Il faut noter que la recherche des composantes présentant des variations significatives n'est pas facile. Il nécessite de surveiller toutes les composantes en plus du problème de détermination de la fenêtre à utiliser pour le calcul de la moyenne.

### 3.3.4 Localisation de défauts multiples

Jusqu'ici nous n'avons considéré que le cas de défaut simple. Il faut noter que les approches de localisation utilisant le principe de structuration et la notion de signatures sont restreintes exclusivement à des cas de défaut unique. Dans le cas de défauts multiples, Staroswiecki et Cassar [90] proposent de générer des signatures supplémentaires pour les défauts multiples en combinant plusieurs signatures de défaut individuelles à l'aide d'un opérateur logique de type OU. Dans le

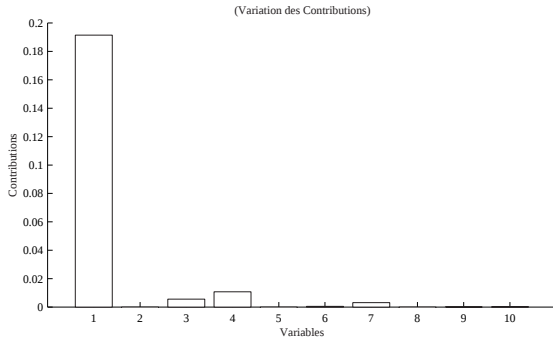


Figure 3.42 – Contributions des différentes variables à l'indice  $D_4$  à l'instant 434, approche des variations

cas d'un système de grande dimension, ceci conduit inévitablement à une explosion combinatoire du nombre de signatures théoriques si l'on considère des possibilités de défauts doubles, triples, etc. Pour réduire le nombre de signature théoriques dans ce cas, Weihua et Sirish [106] proposent de déterminer le nombre maximal de capteurs qui peuvent tomber en panne simultanément et cela par un calcul de la probabilité d'occurrence de défauts multiples.

Or, la méthode de localisation par reconstruction peut être utilisée pour la localisation de défauts multiples en reconstruisant simultanément les variables supposées en défauts.

Pour simplifier le problème, nous allons supposer que l'on veut reconstruire deux variables à l'instant  $k$ ,  $x_i(k)$  et  $x_j(k)$  de  $\mathbf{x}(k)$ . La première chose à faire est de regrouper les deux variables dans un sous-vecteur  $\mathbf{x}_a(k)$  pour mettre  $\mathbf{x}(k)$  sous la forme  $\mathbf{x}(k) = [\mathbf{x}_a(k) \ \mathbf{x}_b(k)]^T$ , où  $\mathbf{x}_b$  représente un sous-vecteur de  $\mathbf{x}$  contenant le reste des variables. En fonction de ces notations l'écriture de l'expression du  $SPE(k)$  est donnée par :

$$SPE(k) = \mathbf{x}^T(k) \tilde{C} \mathbf{x}(k) = \begin{bmatrix} \mathbf{x}_a^T(k) & \mathbf{x}_b^T(k) \end{bmatrix} \begin{pmatrix} \tilde{C}_{aa} & \tilde{C}_{ab} \\ \tilde{C}_{ba} & \tilde{C}_{bb} \end{pmatrix} \begin{bmatrix} \mathbf{x}_a(k) \\ \mathbf{x}_b(k) \end{bmatrix} \quad (3.169)$$

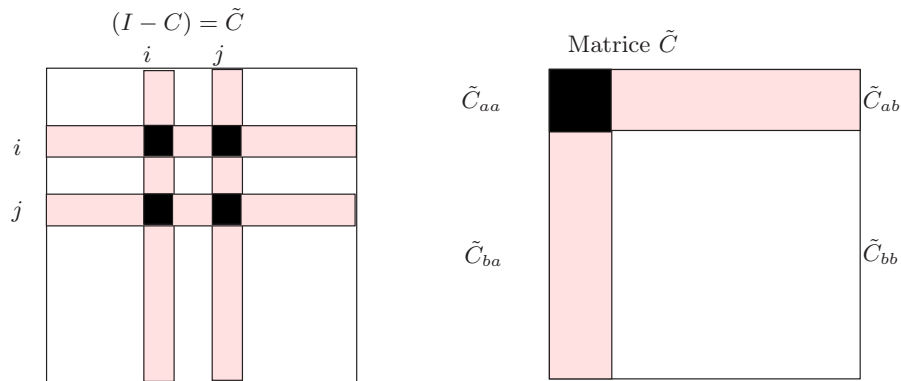


Figure 3.43 – Réorganisation de la matrice  $\tilde{C}$  pour les défauts multiples

Après développement de l'expression du  $SPE$  et en tenant en compte que  $\tilde{C}_{ba} = \tilde{C}_{ab}^T$ , l'expression suivante du  $SPE(k)$  est obtenue.

$$SPE(k) = \mathbf{x}_a^T(k) \tilde{C}_{aa} \mathbf{x}_a(k) + \mathbf{x}_b^T(k) \tilde{C}_{ba} \mathbf{x}_a(k) + \mathbf{x}_a^T(k) \tilde{C}_{ba}^T \mathbf{x}_b(k) + \mathbf{x}_b^T(k) \tilde{C}_{bb} \mathbf{x}_b(k) \quad (3.170)$$

La reconstruction du sous-vecteur  $\mathbf{x}_a$ , qui sera noté  $\mathbf{x}_a^{(r)}(k)$ , est obtenue par minimisation du  $SPE(k)$  par rapport à ce vecteur :

$$\mathbf{x}_a^{(r)}(k) = -\tilde{C}_{aa}^{-1} \tilde{C}_{ba}^T \mathbf{x}_b(k) \quad (3.171)$$

Ce résultat est donné sous la condition que  $\tilde{C}_{aa}$  soit inversible. Dans le cas de défaut simple, nous avons  $\mathbf{x}_a = x_i$ ,  $\mathbf{x}_b(k) = \mathbf{x}^{(i)}(k)$  où  $\mathbf{x}^{(i)}(k)$  est un vecteur contenant toutes les variables sauf la  $i^{eme}$ .  $\tilde{C}_{aa} = 1 - c_{ii}$ ,  $\tilde{C}_{ba}^T = [-c_{1i} \dots -c_{(i-1)i} \quad -c_{(i+1)i} \dots -c_{mi}]$ , et ainsi :

$$\mathbf{x}_i^{(r)}(k) = -\tilde{C}_{aa}^{-1} \tilde{C}_{ba}^T \mathbf{x}_b(k) = \frac{1}{1 - c_{ii}} [c_{1i} \dots c_{(i-1)i} \quad c_{(i+1)i} \dots c_{mi}] \mathbf{x}_b(k) \quad (3.172)$$

Cette dernière équation est la même que celle donnée par l'équation (2.41).

Cette expression peut être exprimée en fonction des directions de défaut comme dans le cas de défaut simple. Ainsi, en posant  $\Xi_i$  comme la matrice des directions du  $i^{eme}$  défaut, qui est supposé affecté plusieurs variables, et  $\Xi_i = \hat{\Xi}_i + \tilde{\Xi}_i$  où  $\hat{\Xi}_i = \hat{C} \Xi_i$  et  $\tilde{\Xi}_i = \tilde{C} \Xi_i$ . Comme dans le cas simple, nous pouvons écrire [18] :

$$\mathbf{x}_i(k) = \left( I - \Xi_i \left( \tilde{\Xi}_i^T \tilde{\Xi}_i \right)^{-1} \tilde{\Xi}_i^T \right) \mathbf{x}(k) \quad (3.173)$$

à condition que la matrice  $(\tilde{\Xi}_i^T \tilde{\Xi}_i)$  soit inversible.

Le principe de localisation dans ce cas reste le même que dans le cas de défaut simple. Toutefois, il faut considéré toutes les combinaisons de défauts possibles. La reconstruction simultanée des capteurs en défauts éliminera l'effet des défauts et l'indice de détection résultant sera en dessous de son seuil de détection. La figure (fig 3.44) présente l'évolution des différents indices obtenus par reconstruction des variables correspondant aux différentes combinaisons de défauts possibles, dans le cas d'un défaut affectant les variables  $x_3$  et  $x_4$  de l'exemple 1. Il est clair que le seul indice qui est en dessous de son seuil de détection, est  $SPE_{34}$  obtenu par reconstruction simultanée des deux variables  $x_3$  et  $x_4$ .

La direction du défaut sera donnée par :

$$\Xi_{3,4} = \begin{pmatrix} 00 \\ 00 \\ 10 \\ 01 \\ 00 \\ 00 \\ 00 \\ 00 \end{pmatrix} \quad (3.174)$$

Ainsi, le vecteur  $\mathbf{x}$  dont les deux variables  $x_3$  et  $x_4$  sont reconstruites sera donnée par :

$$\mathbf{x}_{3,4}(k) = \begin{pmatrix} 1 & 0 & 00 & 0 & 0 & 0 \\ 0 & 1 & 00 & 0 & 0 & 0 \\ 0.33 & 0.33 & 00 & -0.19 & 0.16 & 0.20 \\ -0.19 & -0.19 & 00 & 0.75 & 0.27 & 0.18 \\ 0 & 0 & 00 & 1 & 0 & 0 \\ 0 & 0 & 00 & 0 & 1 & 0 \\ 0 & 0 & 00 & 0 & 0 & 1 \end{pmatrix} \mathbf{x}(k) \quad (3.175)$$

et le vecteur des résidus correspondant est donné par :

$$\mathbf{e}_{3,4}(k) = \begin{pmatrix} 0.6 & -0.33 & 0 & 0 & 0.19 & -0.16 & -0.20 \\ -0.33 & 0.66 & 0 & 0 & 0.19 & -0.16 & -0.20 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0.19 & 0.19 & 0 & 0 & 0.24 & -0.27 & -0.18 \\ -0.16 & -0.16 & 0 & 0 & -0.27 & 0.70 & -0.27 \\ -0.20 & -0.20 & 0 & 0 & -0.18 & -0.27 & 0.72 \end{pmatrix} \mathbf{x}(k) \quad (3.176)$$

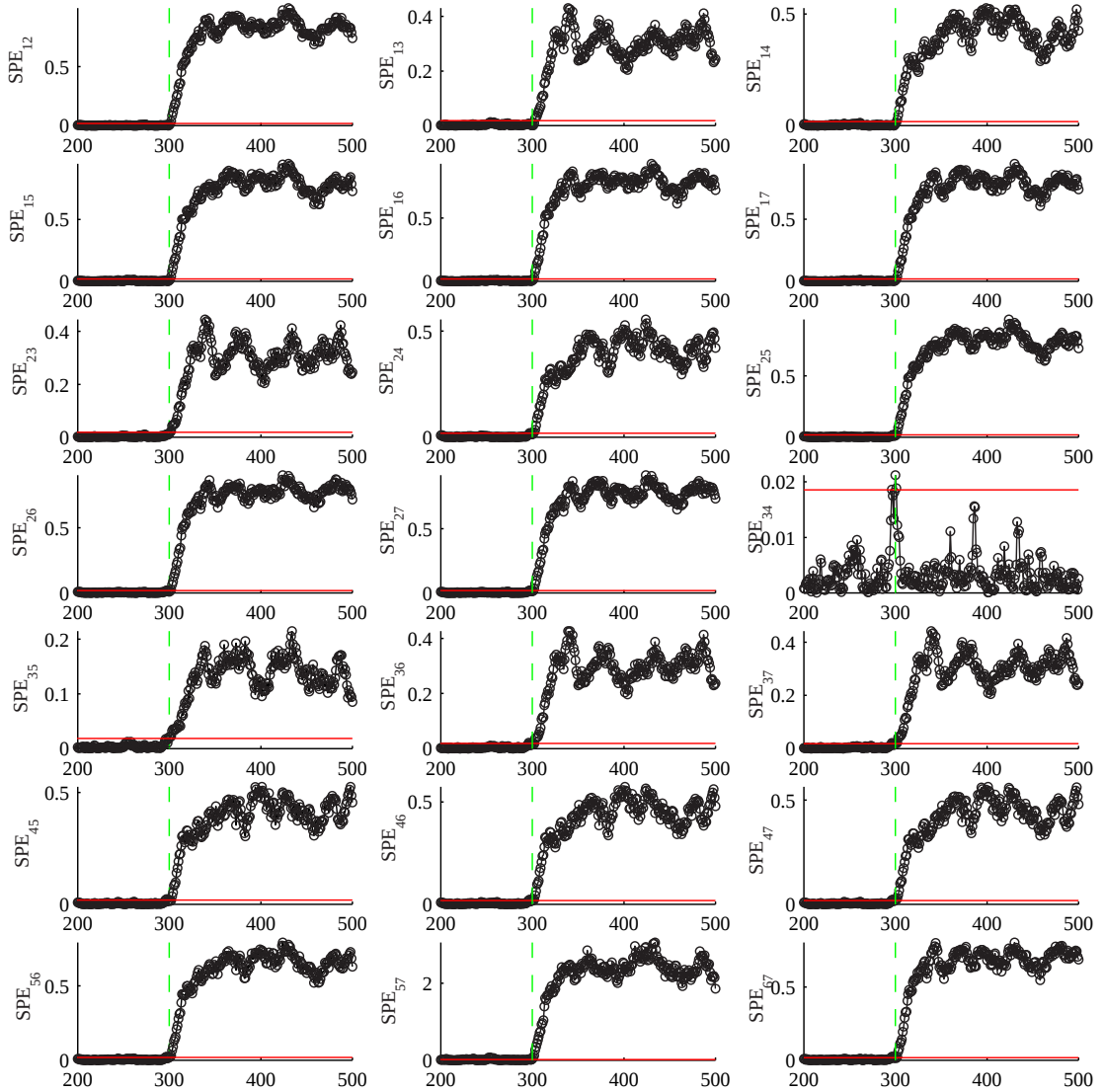


Figure 3.44 – Résultat de la localisation par reconstruction d'un défaut affectant les variables  $x_3$  et  $x_4$ .

Cependant, il faut noter que dans le cas de défauts multiples la méthode de localisation par reconstruction nécessite un nombre très grand de combinaisons à tester. Comme le nombre

de capteurs défaillants n'est pas connu a priori, on commence par la reconstruction d'un seul capteur. Si l'indice de détection révèle toujours la présence d'un défaut pour chaque capteur reconstruit, alors on passe à la reconstruction de deux capteurs en prenant en compte toutes les combinaisons possibles.

Pour cette raison nous proposons d'adapter l'approche de reconstruction à notre indice de détection  $D_i$ . Le nombre de combinaison à tester est considérablement réduit. Pour illustrer l'approche de localisation de défauts multiples avec l'indice  $D_i$ , nous allons considérer l'exemple 2 avec deux défauts affectant les variables  $x_3$  et  $x_9$  avec les mêmes défauts que précédemment entre les instants 300 et 500.

La procédure de détection est lancée et un défaut est détecté sur l'indice  $D_1$ . Aussitôt la procédure de localisation en utilisant l'approche de reconstruction est exécutée et la variable incriminée  $x_9$  est localisée comme le montre la figure (fig 3.45). Une fois la variable localisée, elle est reconstruite et la procédure de détection est relancée. Ainsi les indices de détection  $D_i^{(9)}$  avec la variable  $x_9$  reconstruite sont calculés. L'indice  $D_3^{(9)}$  indique la présence d'un défaut. Pour localiser la deuxième variable on n'a plus besoin de tester toutes les combinaisons deux à deux mais juste celles qui sont liées à  $x_9$ . De plus en analysant la matrice des derniers vecteurs propres (tab 3.9) et plus particulièrement les vecteurs  $p_8$ ,  $p_9$  et  $p_{10}$  qui interviennent dans le calcul de l'indice  $D_4$ , on constate que les variables  $x_1$ ,  $x_4$  et  $x_5$  ont des coefficients très faibles et donc ne peuvent être en cause du défaut détecté. Ainsi, ces variables ne sont pas reconstruites. Les variables à reconstruire sont  $x_2, x_3, x_6, x_7, x_8, x_{10}$  en combinaison avec la variable  $x_9$ . Le résultat de l'application de cette procédure est illustré sur la figure (fig 3.46) et la deuxième variable est localisée.

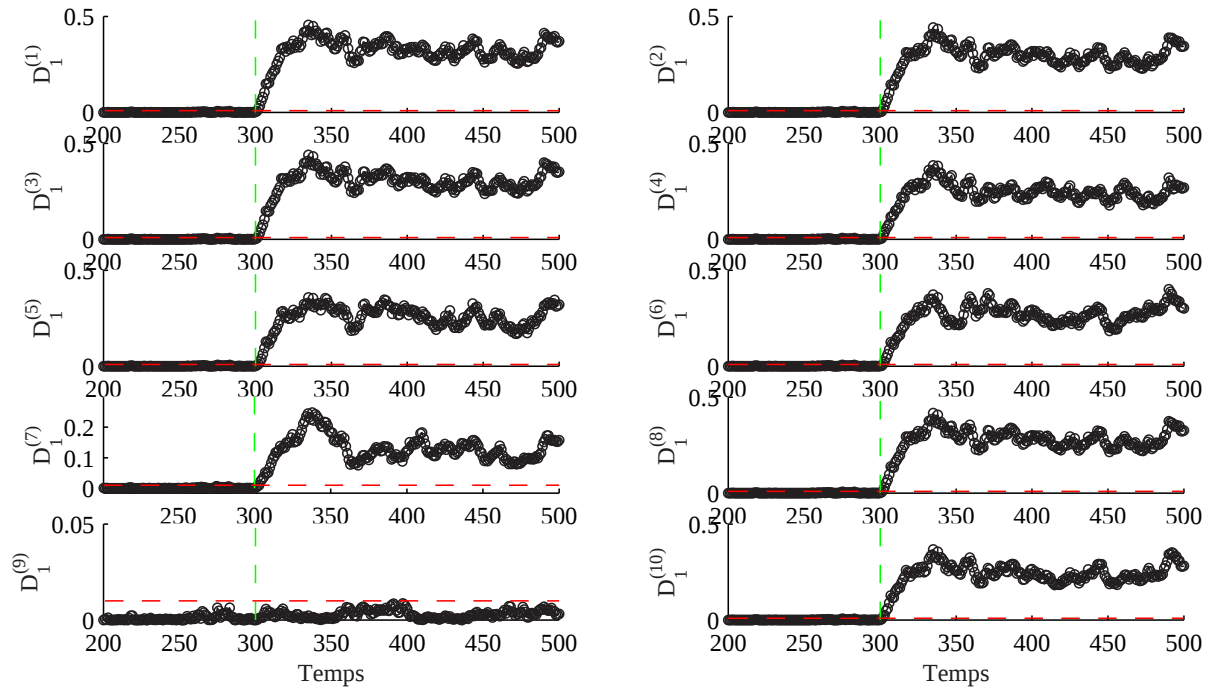


Figure 3.45 – Résultat de la localisation par reconstruction de la variable  $x_9$  avec un défaut affectant les variables  $x_3$  et  $x_9$ .

Cependant, cette procédure n'est pas automatique. Ainsi, il est intéressant d'automatiser cette

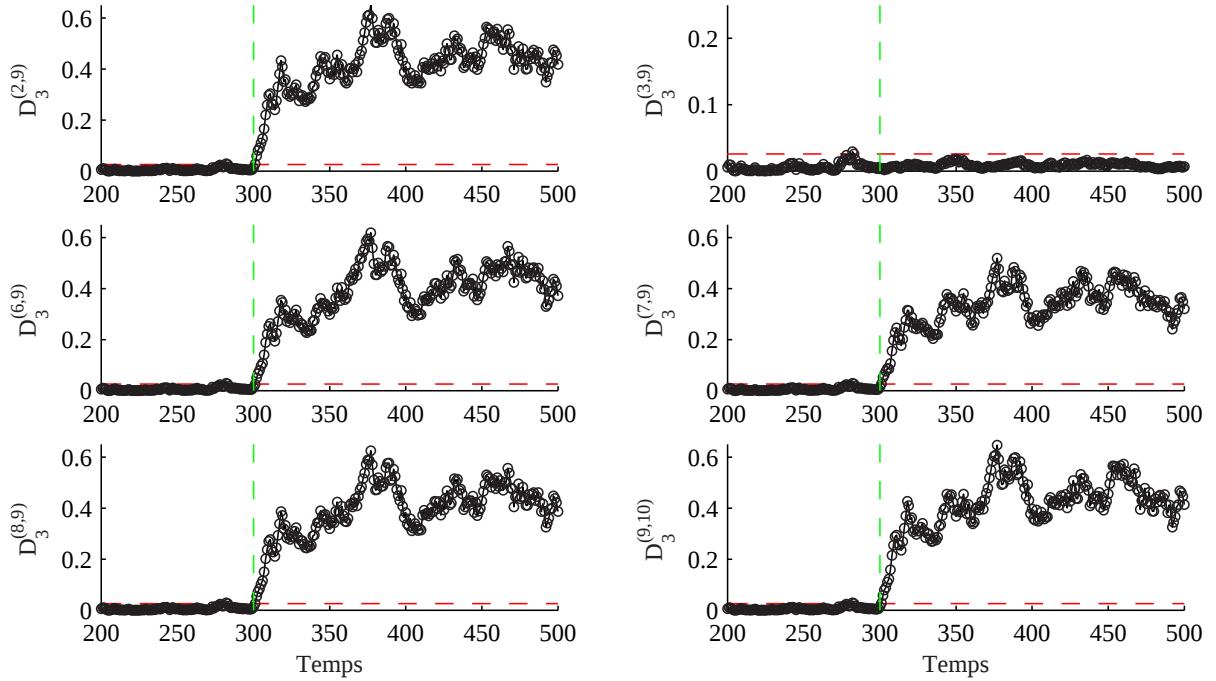


Figure 3.46 – Résultat de la localisation par reconstruction de la variable  $x_3$  avec un défaut affectant les variables  $x_3$  et  $x_9$ .

procédure de localisation de défauts multiples en tenant compte des sensibilités des coefficients intervenants dans le calcul de l'indice de détection  $D_i$ .

### 3.4 Identification des caractéristiques du défaut

Comme nous l'avons déjà vu dans le chapitre précédent, il y a toujours une partie de la variation des mesures qui ne peut être reconstruite en utilisant les autres capteurs ; cette partie est l'erreur de reconstruction. Dans le cas d'un défaut, le vecteur  $\mathbf{x}$  est reconstruit en utilisant la direction du défaut  $\xi_i$  et son amplitude  $d$ . Ainsi  $\mathbf{x}(k)$  peut être représenté par l'expression :

$$\mathbf{x}(k) = \mathbf{x}^*(k) + \xi_i d(k) \quad (3.177)$$

où  $\mathbf{x}$  est alors le vecteur de mesures affecté par le défaut  $d$  dans la direction  $\xi_i$ .  $\mathbf{x}^*$  est le vecteur de mesure supposé sans défaut. La reconstruction des mesures pour des capteurs en défaut consiste à trouver une estimation de  $\mathbf{x}^*$ , c'est-à-dire corriger au mieux l'effet du défaut. Ainsi le vecteur dont la  $i^{eme}$  variable est reconstruite est donné par :

$$\mathbf{x}_i = \mathbf{x} - d_i \xi_i \quad (3.178)$$

où  $\mathbf{x}_i$  est le vecteur de mesure reconstruit, et  $d_i$  est l'estimation de l'amplitude du défaut  $d$  dans la direction  $\xi_i$ .

En d'autres termes, il faut trouver  $d_i$  tel que l'erreur entre la quantité reconstruite (3.178) et son estimation par le modèle ACP soit minimisée :

$$d_i = \arg \min_{d_i} \|\mathbf{x}_i - \hat{\mathbf{x}}_i\|^2 = \arg \min_{d_i} \|\tilde{\mathbf{x}}_i\|^2 = \arg \min_{d_i} \|\tilde{\mathbf{x}} - d_i \tilde{\xi}_i\|^2 \quad (3.179)$$

où  $\tilde{\xi}_i$  est la projection de  $\xi_i$  dans le sous-espace des résidus,  $\tilde{\xi}_i = (I - \hat{C})\xi_i$  et  $\hat{\mathbf{x}}_i$  représente l'estimation de  $\mathbf{x}_i$  fournie par le modèle ACP. La solution du problème de minimisation de l'équation (3.179) est donnée par :

$$d_i = \frac{\tilde{\xi}_i^T \tilde{\mathbf{x}}}{\tilde{\xi}_i^T \tilde{\xi}_i} = \frac{\tilde{\xi}_i^T \mathbf{x}}{\tilde{\xi}_i^T \xi_i} \quad (3.180)$$

### Conditions nécessaires et suffisantes pour l'estimation du défaut

A partir de l'équation (3.180),  $d_i$  peut être calculer si  $\|\tilde{\xi}_i\| \neq 0$ . Donc, pour estimer le défaut, on retrouve la même condition que pour la reconstruction.

## 3.5 Conclusion

Ce chapitre a été consacré à la présentation des techniques de détection et de localisation de défaut dans le cadre d'une procédure de diagnostic utilisant l'analyse en composantes principales linéaires.

Avec l'ACP, on peut générer des résidus qui sont équivalents à ceux que l'on peut obtenir par l'espace de parité.

Plusieurs indices de détection ont été présentés dans le cas des approches par estimation d'état et par estimation paramétrique. Les indices de détection les plus utilisés dans le cadre de l'ACP sont les statistiques  $T^2$  et l'erreur quadratique d'estimation  $SPE$ .

Cependant, le  $SPE$  est sensible aux erreurs de modélisation [31] car c'est un test global qui cumule les erreurs de chacun des résidus et la statistique  $T^2$  telle qu'elle a été définie avec l'ACP ne peut être utilisée pour la détection car elle est calculée à partir des premières composantes principales et ces dernières ne sont pas des résidus. Pour cette raison nous avons proposé un nouvel indice de détection basé sur les dernières composantes principales, permettant de s'affranchir des inconvénients des indices existants et d'améliorer la qualité de la détection [31]. Cet indice s'exprime comme une somme des carrés des dernières composantes principales dans différents sous-espaces de l'espace résiduel. Ainsi, il représente un  $SPE$  calculé dans des sous-espaces plus petits que l'espace résiduel et donc il présente l'avantage d'être plus sensible aux défauts que les autres indices.

Différentes approches de localisation ont été également présentées. Dans un premier temps nous avons présenté le principe des méthodes de localisation avec structuration des résidus dans le cas de l'ACP et les différentes approches proposées dans la littérature. Nous avons mis en évidence les limites d'utilisation de ces approches. Il faut noter que ces approches cherchent une matrice de signatures théorique fortement localisante sans tenir compte de la sensibilité des résidus primaires aux différents défauts.

La deuxième approche présentée dans ce chapitre concerne l'approche utilisant un banc de modèles comme dans les approches utilisant des observateurs (GOS). Dans cette approche on retrouve trois méthodes. La méthode utilisant les ACP partielles, la méthode utilisant le principe de reconstruction et la méthode utilisant le principe d'élimination. Cependant, dans le cas de la méthode des ACP partielles on ne sait pas s'il est possible d'élaborer un modèle ACP réduit a priori d'autant plus quelle utilise le  $SPE$  comme indice de détection. La deuxième méthode présentée utilise le principe de reconstruction qui consiste à suspecter qu'un capteur est défaillant

et à reconstruire la valeur de sa mesure en se basant sur le modèle ACP et les mesures des autres capteurs. Comme cette méthode utilise initialement l'indice *SPE* nous l'avons adaptée à l'indice de détection  $D_i$ . De cette manière, la reconstruction se fait avec le modèle qui donne la meilleurs reconstruction alors que la détection se fait avec un autre modèle (en utilisant l'indice  $D_i$ ) permettant d'être de plus sensible aux défauts. La procédure de localisation proposée permet de localiser des défauts simples et peut être utilisée pour la localisation de défauts multiples. Nous avons présenté le principe d'utilisation de cette approche pour les défauts multiples mais il reste à automatiser la procédure proposée.

De même pour la méthode utilisant le principe d'élimination, cette approche a été adaptée à l'indice de détection proposé et donne de bons résultats comme nous l'avons vu sur les exemples de simulation.

La troisième approche qui a été abordée dans ce chapitre concerne l'approche par calcul des contributions des variables à l'indice de détection. Cette approche est largement utilisée dans la littérature pour la localisation de défauts par l'analyse en composantes principales. A partir de l'expression de l'indice  $D_i$ , nous avons proposé deux définitions des contributions des variables à l'indice proposé. La première définition est basée sur le fait que cet indice est un *SPE* particulier et de ce fait on peut exploiter la définition des contributions dans le cas du *SPE*. La deuxième définition vient du fait que l'indice proposé peut être calculé également à partir des dernières composantes principales. Ainsi, un calcul des contributions aux dernières composantes ayant subi une variation significative peut être utilisé pour cette définition. Dans les deux définitions, la variable ayant la plus grande contribution est considérée comme la variable incriminée.

Les deux exemples de simulation du chapitre 2 ont été utilisés pour illustrer les différentes approches présentées dans ce chapitre.



## 4

Analyse en composantes principales non linéaires  
(ACPNL)

## Sommaire

|            |                                                                       |            |
|------------|-----------------------------------------------------------------------|------------|
| <b>4.1</b> | <b>Introduction</b>                                                   | <b>101</b> |
| <b>4.2</b> | <b>Analyse en composantes principales non linéaires</b>               | <b>103</b> |
| <b>4.3</b> | <b>Méthode des Courbes Principales</b>                                | <b>104</b> |
| 4.3.1      | Algorithme de calcul des courbes principales de Hastie                | 107        |
| 4.3.2      | Algorithme de calcul des courbes principales de Verbeek               | 107        |
| <b>4.4</b> | <b>Approches neuronales de l'ACPNL</b>                                | <b>110</b> |
| 4.4.1      | ACPNL par réseau à cinq couches                                       | 111        |
| 4.4.2      | ACPNL par optimisation des entrées du réseau (Input Training network) | 116        |
| 4.4.3      | Réseaux à trois couches et courbes principales                        | 117        |
| 4.4.4      | Réseaux de fonctions à base radiale (RBF)                             | 119        |
| <b>4.5</b> | <b>Conclusion</b>                                                     | <b>129</b> |

## 4.1 Introduction

L'ACP linéaire, comme elle a été introduite dans le premier chapitre, est une des méthodes d'analyse de données les plus connues. Elle consiste à chercher un sous-espace de dimension plus petite que l'espace de départ et à y projeter les données étudiées, en perdant un minimum d'information. Le résultat ainsi obtenu est une représentation des données avec réduction de dimension.

Pour réduire les calculs, dans le cas où la matrice de corrélation est de grande dimension, des implantations sous forme de réseaux de neurones de l'ACP linéaire ont été proposées ([72], [73], [84], [83], [59]). Les approches neuronales de l'analyse en composantes principales linéaires se distinguent généralement par deux critères d'apprentissage optimisés qui sont d'ailleurs équivalents : maximisation des variances de projection des données [41], et minimisation de l'erreur quadratique d'estimation des données [75].

Malheureusement, comme c'est une opération de projection linéaire, seules les dépendances linéaires ou quasi-linéaires entre les variables peuvent être révélées. Si les données à traiter présentent des comportements non linéaires, l'ACP est incapable de trouver la représentation compacte décrivant ces données. Ainsi, l'extension de l'analyse en composantes principales pour traiter les problèmes non linéaires a été développée [28, 38, 56, 12, 92, 104, 70, 49]. Hastie [38, 39] propose une approche pour une généralisation de l'ACP dans le cas non linéaire basée sur le principe des courbes principales. Une courbe principale est une courbe lisse minimisant la distance entre tous les points de données et leurs projections sur cette courbe. Toutefois, cette approche est non paramétrique (pas de modèle de représentation) et ne peut être utilisée pour la surveillance. De plus elle ne permet de calculer que des composantes principales non linéaires unidimensionnelles. Indépendamment des travaux de Hastie, Kramer [56] propose une analyse en composantes principales non linéaires (ACPNL) en utilisant un réseau de neurones à cinq couches dont les poids sont calculés par apprentissage en minimisant l'erreur quadratique entre les entrées et les sorties du réseau. Vue la complexité d'un tel réseau, Dong [12] propose une ACPNL combinant les réseaux de neurones et l'approche des courbes principales, deux sous-réseaux à trois couches étant utilisés dans cette approche. Les composantes principales sont obtenues à partir de l'algorithme des courbes principales de Hastie [39, 96] et le problème est transformé en un problème de régression non-linéaire. Tan et Mavrovouniotis [92] ont proposé une ACPNL obtenue en utilisant un réseau de neurones à trois couches et dont l'apprentissage est effectué en minimisant à la fois les poids et les entrées du réseau. Cependant, pour de tels réseaux, l'apprentissage est très lourd et demande beaucoup de temps de calcul en plus des problèmes d'initialisation et de convergence. Pour cette raison, Webb [104] a proposé une approche pour l'analyse en composantes principales en utilisant deux réseaux de fonctions de bases radiales (RBF) à trois couches en cascade. Le même principe a été adopté par Wilson [111]. Cependant l'apprentissage d'un tel réseau reste très complexe.

Comme, la plupart des approches utilisent les réseau de neurones MLP (multi-layer perceptrons) pour l'obtention du modèle ACPNL, on rencontre souvent des problèmes d'optimisation non linéaires comme les problème de convergence et d'initialisation de ce type de réseaux. Pour cette raison et en combinant les courbes principales et les réseaux RBF, nous proposons une approche pour l'ACPNL avec deux réseaux à trois couches en cascade ou le problème d'apprentissage se ramène a un problème de regression linéaire par rapport aux poids de la couche de sortie.

De plus, comme dans le cas linéaire, nous proposons un algorithme permettant de déterminer le nombre de composantes non-linéaires à retenir dans le modèle ACPNL. Cette algorithme est une extension du principe de la variance non reconstruite que nous avons présenté dans le cas linéaire.

Dans ce chapitre nous allons présenter la généralisation de l'ACP linéaire dans le cas non linéaire. Dans un premier temps, le principe de l'analyse en composantes principales non linéaires est présenté ainsi que les principales approches pour le calcul du modèle ACPNL.

Il faut noter que dans notre travail, on ne s'intéresse pas à la recherche de la structure des réseaux étudiés. Cependant nous nous sommes intéressés plus particulièrement au nombre de composantes à retenir dans le modèle ACPNL. Ainsi, nous proposons une approche permettant de déterminer le nombre de composantes dans le cas non linéaire basée sur l'extension du principe de la variance non reconstruite dans le cas non linéaire.

Tout au long de ce chapitre, un exemple non-linéaire simple sera présenté pour illustrer les méthodes présentées. L'exemple choisi est décrit par les équations suivantes :

$$\begin{aligned}
x_1 &= u^2 + 0.3 \sin(2\pi u) + \varepsilon_1 \\
x_2 &= u + \varepsilon_2 \\
x_3 &= u^3 + u + 1 + \varepsilon_3
\end{aligned} \tag{4.1}$$

où  $u \in [-1, 1]$  est une variable aléatoire uniforme et  $\varepsilon_i$  un bruit aléatoire uniformément distribué dans l'intervalle  $[-0.1, 0.1]$ .

## 4.2 Analyse en composantes principales non linéaires

L'ACP non linéaire est une extension de l'ACP linéaire. Tandis que l'ACP cherche à identifier les relations linéaires entre les variables du processus, l'objectif de l'ACP non linéaire est d'extraire à la fois les relations linéaires et non linéaires. Cette généralisation est effectuée par une projection des données du processus sur des courbes ou des surfaces (fig 4.2) au lieu des droites ou des plans (fig 4.1).

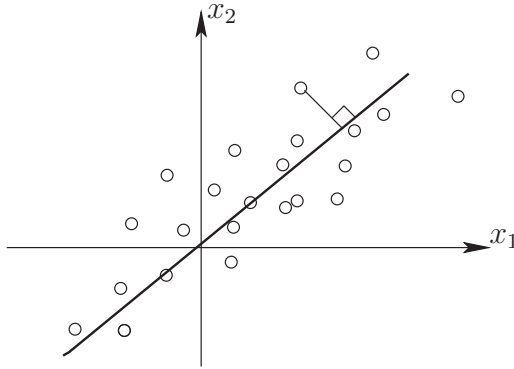


Figure 4.1 – ACP linéaire.

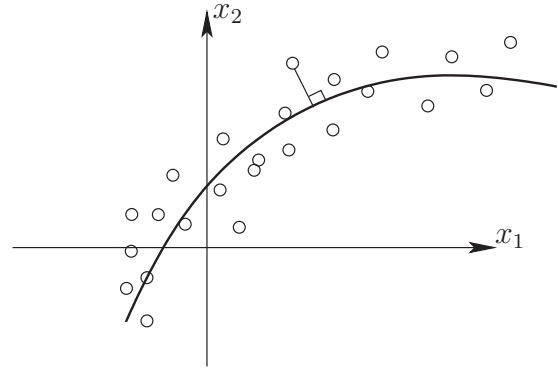


Figure 4.2 – ACP non linéaire.

Pour mieux comprendre le problème et pouvoir faire le lien avec le modèle linéaire, nous allons considérer la figure (fig 4.3) qui représente le principe du modèle ACP général, que se soit le modèle linéaire ou non linéaire. Le modèle global est composé de deux sous-modèles, un sous-modèle de compression projette des données de dimension  $m$  vers l'espace des composantes principales de dimension  $\ell$  et le deuxième sous-modèle effectue l'opération inverse, à savoir une projection de  $\mathbb{R}^\ell$  vers  $\mathbb{R}^m$ . Ainsi dans le cas linéaire nous avons les deux sous-modèles qui sont donnés par les deux matrices orthogonales des vecteurs propres de la matrice de corrélation des données :  $\hat{P}$  et  $\hat{P}^{-1} = \hat{P}^T$  et le modèle global est donné par la matrice de projection  $\hat{C}$  définie par  $\hat{C} = \hat{P}\hat{P}^T$ . Dans le cas non linéaire, le but est de chercher deux fonctions non linéaires  $\mathcal{F}$  et  $\mathcal{G}$ .  $\mathcal{G}$  représente le modèle non linéaire de compression qui permet de calculer les composantes principales non linéaires à partir des données et  $\mathcal{F}$  représente le modèle non linéaire de décompression permettant l'estimation des variables originelles à partir des composantes principales non linéaires données par le modèle de compression.

Ainsi, la matrice de données  $X$  peut être représentée par une estimation  $\hat{X}$ , donnée par le modèle ACP, plus une erreur d'estimation  $E$  :

$$X = \hat{X} + E = \mathcal{F}(T) + E \tag{4.2}$$

où  $T = [t_1, \dots, t_\ell] \in \mathbb{R}^{N \times \ell}$  est la matrice des composantes principales non linéaires qui est donnée par :

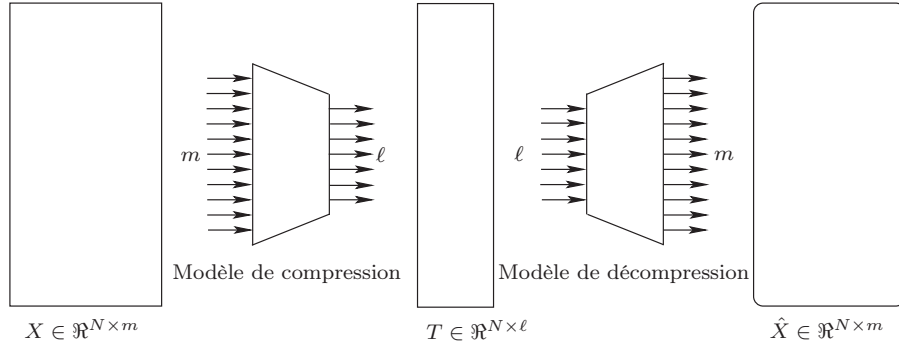


Figure 4.3 – Principe de la modélisation par l’analyse en composantes principales

$$T = \mathcal{G}(X) \quad (4.3)$$

A partir de cette équation le modèle de l’ACP non linéaire (ACPNL) est représenté par la fonction  $\mathcal{F}(\mathcal{G}(\cdot))$  et l’estimation de  $X$ , notée  $\hat{X}$ , est donnée par :

$$\hat{X} = \mathcal{F}(\mathcal{G}(X)) \quad (4.4)$$

Gnanadesikan [28] propose l’un des premiers algorithmes pour l’analyse en composantes principales non linéaires : l’analyse en composantes principales généralisées. Pour explorer la structure non linéaire des données, Gnanadesikan propose d’ajouter des termes quadratiques et peut-être des termes avec un ordre plus grand à la liste des variables et d’appliquer l’ACP linéaire à cet ensemble étendu de variables. Par exemple si nous avons deux variables  $x_1$  et  $x_2$ , l’ACP linéaire est appliquée sur les cinq variables  $x_1, x_2, x_1^2, x_2^2$  et  $x_1x_2$ . Cette approche n’a pas connue d’application réelle, parce que les transformations à appliquer aux variables ne sont pas évidentes. De plus les relations linéaires que nous allons chercher ne sont pas des relations exactes mais des relations entre des variables entachées de bruit de mesure car les transformations sont effectuées sur les variables mesurées.

Dans la suite nous allons présenter les différentes approches pour l’extraction des composantes principales non linéaires. Dans un premier temps nous présenterons le principe des courbes principales, ensuite seront présentées les approches neuronales pour l’ACPNL utilisant soit les réseaux de neurones MLP soit les réseaux de fonctions à base radiale.

### 4.3 Méthode des Courbes Principales

Hastie [38] a proposé une approche pour l’ACP non linéaire basée sur la méthode des courbes principales. Elle est basée sur la généralisation des propriétés statistiques de l’ACP linéaire. Considérant un vecteur de données de dimension  $m$ ,  $\mathbf{x} = [x_1, \dots, x_m]^T$ . La première composante principale de  $\mathbf{x}$  est une droite minimisant la distance euclidienne entre elle et le vecteur  $\mathbf{x}$ . Ainsi, la première composante principale est la meilleure approximation unidimensionnelle de  $\mathbf{x}$ , et la projection de  $\mathbf{x}$  sur cette droite donne la meilleure représentation linéaire des données. Ainsi, voici quelques définitions de la géométrie des courbes.

**Définition 4.3.1** Une courbe dans un espace Euclidien est une fonction continue  $\mathcal{F} : I \rightarrow \mathbb{R}^m$ , où  $I = [a, b]$  est un interval fermé de  $\mathbb{R}$ . La courbe  $\mathcal{F}$  peut être considérée comme un vecteur

de fonctions de dimension  $m$  de la variable  $t$ ,  $\mathcal{F}(t) = (f_1(t), \dots, f_m(t))$ , où  $f_1(t), \dots, f_m(t)$  sont appelées fonctions coordonnées.

**Définition 4.3.2 (Longueur d'une courbe)** La longueur de la courbe  $\mathcal{F}$  le long d'un intervalle  $[\alpha, \beta] \subset [a, b]$ , notée par  $l(\mathcal{F}, \alpha, \beta)$ , est définie par :

$$l(\mathcal{F}, \alpha, \beta) = \sup \sum_{i=1}^N \|\mathcal{F}(t_i) - \mathcal{F}(t_{i-1})\| \quad (4.5)$$

où  $\alpha = t_0 \leq t_1 < \dots \leq t_N = \beta$ ,  $N \geq 1$ . La longueur de  $\mathcal{F}$  le long de tout son domaine  $[a, b]$  est notée par  $l(\mathcal{F})$ .

**Définition 4.3.3 (Distance entre un point et une courbe)** Soit  $\mathcal{F}(t) = (f_1(t), \dots, f_m(t))$  une courbe lisse (dont les  $f_i$  sont infiniment dérivables) dans  $\mathbb{R}^m$  paramétrée par  $t \in \mathbb{R}$ , et pour  $\mathbf{x} \in \mathbb{R}^m$  soit  $t_{\mathcal{F}}(\mathbf{x})$  la valeur du paramètre  $t$  pour laquelle la distance entre  $\mathbf{x}$  et  $\mathcal{F}(t)$  est minimisée. Plus formellement, l'indice de projection  $t_{\mathcal{F}}(\mathbf{x})$  est défini par l'équation :

$$t_{\mathcal{F}}(\mathbf{x}) = \sup \left\{ t : \|\mathbf{x} - \mathcal{F}(t)\| = \inf_{\tau} \|\mathbf{x} - \mathcal{F}(\tau)\| \right\} \quad (4.6)$$

où  $\|\cdot\|$  est la norme euclidienne.

**Définition 4.3.4 (Propriétés géométriques des courbes)** Soit  $\mathcal{F} : [a, b] \rightarrow \mathbb{R}^m$  une courbe lisse  $\mathcal{F} = (f_1, \dots, f_m)$  dont la dérivée est donnée par :

$$\mathcal{F}'(t) = \left( \frac{\partial f_1(t)}{\partial t}, \dots, \frac{\partial f_m(t)}{\partial t} \right) \quad (4.7)$$

La  $\mathcal{F}'(t)$  est la tangente à la courbe au point  $t$  et pour une courbe paramétrée par sa longueur d'arc  $\|\mathcal{F}'(t)\| = 1$  [58]. Notons que pour une courbe lisse  $\mathcal{F} : [a, b] \rightarrow \mathbb{R}^m$ , la longueur de la courbe (4.5) le long de l'intervalle  $[\alpha, \beta] \subset [a, b]$  peut être définie comme :

$$l(\mathcal{F}, \alpha, \beta) = \int_{\alpha}^{\beta} \|\mathcal{F}'(t)\| dt \quad (4.8)$$

La longueur d'arc est indépendante du choix de la représentation paramétrique. Ainsi, les mesures géométriques basées sur la longueur d'arc sont invariantes par rapport à la paramétrisation.

Une courbe principale est définie comme une courbe auto-consistante, où la propriété d'auto-consistance peut s'interpréter par le fait que chaque point de la courbe  $\mathcal{F}$  est la moyenne de tous les points qui sont projetés sur elle. Ainsi, les courbes principales sont des courbes lisses auto-consistantes qui passent au milieu du nuage de point de dimension  $m$  et donne un résumé unidimensionnel non linéaire des données [39].

Généralement la courbe est paramétrée par sa longueur d'arc, c'est-à-dire que chaque point sur la courbe peut être décrit par sa distance le long de la courbe à partir de l'origine. La calcul de la longueur d'arc de la courbe  $\mathcal{F}$  entre  $t_0$  et  $t_1$  est donnée par :

$$l(\mathcal{F}, t_0, t_1) = \int_{t_0}^{t_1} \sqrt{\sum_{j=1}^m \left( \frac{\partial \mathcal{F}_j}{\partial t} \right)^2}_{(t=z)} dz \quad (4.9)$$

Comme la définition de la courbe n'est pas unique, il y a différentes fonctions  $\mathcal{F}$  qui définissent la même courbe, mais avec une paramétrisation différente. Ainsi, une propriété additionnelle pour l'unicité de la fonction est donnée par :

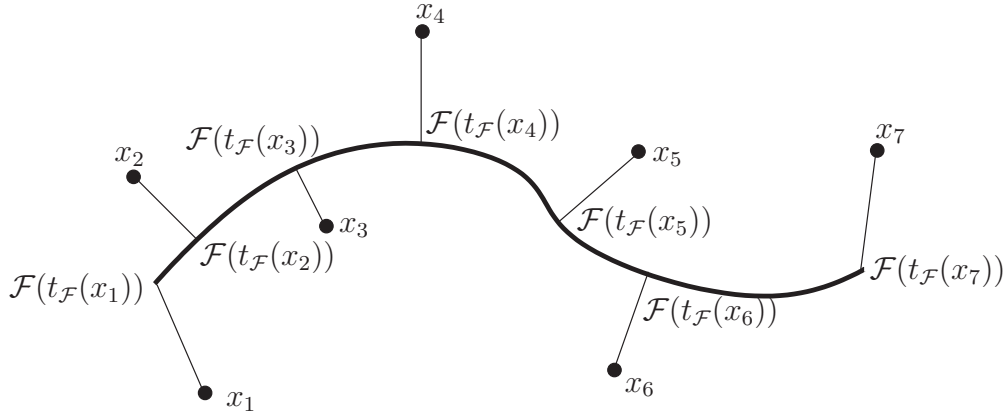


Figure 4.4 – Projection des points sur la courbe.

$$\sum_{j=1}^m \left( \frac{\partial \mathcal{F}_j}{\partial t} \right)^2 = 1 \quad (4.10)$$

Dans le cas linéaire nous avons la même propriété, à savoir que dans le cas d'une seule composante  $\hat{\mathbf{x}} = \mathcal{F}(\mathbf{x})$  où  $\mathcal{F}(\mathbf{x}) = p\mathbf{t}$ . Ainsi,  $\sum_{j=1}^m \left( \frac{\partial \mathcal{F}_j}{\partial t} \right)^2 = \|p\| = 1$ .

Il faut noter que la courbe  $\mathcal{F}$  n'a pas une forme paramétrique fixe et peut être représentée sous forme polygonale avec les sommets  $\mathcal{F}(t_1), \dots, \mathcal{F}(t_N)$ .

Avant de présenter l'algorithme de Hastie, nous allons présenter un algorithme pour le calcul des composantes principales linéaires. L'algorithme que nous allons introduire ici pour le calcul des composantes principales linéaires a été récemment proposé par Roweis [82]. La raison pour laquelle nous le présentons ici est qu'il y a une analogie entre cet algorithme qui calcule les composantes principales et l'algorithme de Hastie pour le calcul des courbes principales. L'idée de base consiste à choisir une droite arbitraire, et de projeter les points de données sur cette droite. Ainsi, nous avons les indices de projection, et on cherche la nouvelle droite qui minimise une certaine distance entre les points de données et cette droite. Une fois la nouvelle droite calculée, répéter les deux étapes précédentes jusqu'à convergence. Formellement, soit  $\mathcal{F}^{(j)}$  la droite obtenue à la  $j^{\text{eme}}$  itération et  $t^{(j)} = [t_1^{(j)}, \dots, t_N^{(j)}]^T = [\mathbf{x}_1^T p^{(j)}, \dots, \mathbf{x}_N^T p^{(j)}]^T$  le vecteur des indices de projection des données sur  $\mathcal{F}^{(j)}$ . La fonction de distance de  $\mathcal{F}(t) = tp$  en supposant que le vecteur  $t^{(j)}$  fixe, est définie par :

$$\begin{aligned} \Delta_N(\mathcal{F}|t^{(j)}) &= \sum_{i=1}^N \|\mathbf{x}_i - t_i^{(j)} p\|^2 \\ &= \sum_{i=1}^N \|\mathbf{x}_i\|^2 + \|p\|^2 \sum_{i=1}^N (t_i^{(j)})^2 - 2p^T \sum_{i=1}^N t_i^{(j)} \mathbf{x}_i \end{aligned} \quad (4.11)$$

Ainsi, pour trouver la droite optimale  $\mathcal{F}^{(j+1)}$ , on doit minimiser (4.11) sous la contrainte  $\|p\| = 1$ .

$$p^{(j+1)} = \arg \min_{\|p\|=1} \Delta_N(\mathcal{F}|t^{(j)}) = \frac{\sum_{i=1}^N t_i^{(j)} \mathbf{x}_i}{\left\| \sum_{i=1}^N t_i^{(j)} \mathbf{x}_i \right\|} \quad (4.12)$$

et ainsi  $\mathcal{F}^{(j+1)}(t) = tp^{(j+1)}$ .

L'algorithme est le suivant :

1. poser  $\mathcal{F}^{(0)}(t) = tp^{(0)}$  une droite arbitraire,  $j = 0$ ,
2. projection :  $t^{(j)} = [t_1^{(j)}, \dots, t_N^{(j)}]^T = [\mathbf{x}_1^T p^{(j)}, \dots, \mathbf{x}_N^T p^{(j)}]^T$ ,
3. estimation :  $p^{(j+1)} = \frac{\sum_{i=1}^N t_i^{(j)} \mathbf{x}_i}{\left\| \sum_{i=1}^N t_i^{(j)} \mathbf{x}_i \right\|}$  et  $\mathcal{F}^{(j+1)}(t) = tp^{(j+1)}$ ,
4. arrêter si  $\left(1 - \frac{\Delta_N(\mathcal{F}^{(j+1)})}{\Delta_N(\mathcal{F}^{(j)})}\right)$  est inférieure à un certain seuil. Sinon,  $j = j + 1$  et aller à l'étape 2.

#### 4.3.1 Algorithme de calcul des courbes principales de Hastie

L'algorithme des courbes principales de Hastie [39] est un algorithme itératif, nécessitant une étape d'estimation et une étape de projection pour chaque itération.

##### 4.3.1.1 Etape de projection

A l'itération  $(j + 1)$ , en ayant la courbe principale courante décrite par  $\mathcal{F}^{(j)}(t_i)$ , on trouve les nouveaux indices de projection  $t_i^{(j+1)} = t_{\mathcal{F}^{(j)}}(\mathbf{x}_i)$  :

1. soit  $d_{ik}$  la distance entre  $\mathbf{x}_i$  et le point paramétré par  $t_k \in [t_k^{(j)}, t_{k+1}^{(j)}]$  sur chaque segment joignant  $\mathcal{F}^{(j)}(t_k^{(j)})$  et  $\mathcal{F}^{(j)}(t_{k+1}^{(j)})$ ,
2. poser  $t_i^{(j+1)} = t_k$  si  $d_{ik} = \min_k(d_{ik})$ ,
3. trouver l'interpolation de  $\mathcal{F}^{(j)}$  correspondant à  $t_i^{(j+1)}$ ,
4. remplacer  $t_i^{(j+1)}$  par la longueur d'arc entre  $\mathcal{F}^{(j)}(t_1)$  et  $\mathcal{F}^{(j)}(t_i)$ . A savoir  $t_i^{(j+1)} = t_{i-1}^{(j+1)} + \left\| \mathcal{F}^{(j)}(t_i) - \mathcal{F}^{(j)}(t_{i-1}) \right\|$  et  $t_1^{(j+1)} = 0$ .

##### 4.3.1.2 Etape d'estimation

A l'itération  $(j + 1)$ , ayant les indices de projection  $t_i^{(j+1)}$ , on calcule  $f^{(j+1)}(t_i^{(j+1)})$  comme la solution du problème de minimisation de la fonction coût suivante pour chaque fonction coordonnée de  $\mathcal{F}$  :

$$D^2(f) = \sum_{i=1}^N \left( x_i - f^{(j)}(t_i^{(j+1)}) \right)^2 + \lambda \int \left( \frac{\partial^2 f^{(j)}}{\partial t^2} \right)^2 dt \quad (4.13)$$

où  $\lambda$  est un paramètre de régularisation et la recherche de  $f$  se fait par des méthodes de régression non paramétrique.

#### 4.3.2 Algorithme de calcul des courbes principales de Verbeek

Beaucoup d'autres auteurs se sont penchés sur le problème des courbes principales [96, 9, 52, 101]. Nous allons présenter brièvement l'approche, pour le calcul des courbes principales, développée par Verbeek [101]. L'algorithme proposé exploite l'algorithme des k-lines où, à chaque

centre du nuage de points, on cherche à tracer une droite ou plus exactement un segment de droite pour approximer la courbe  $\mathcal{F}$ .

Verbeek utilise des modèles locaux unidimensionnels représentant ainsi des segments de droite et le raccordement entre les segments adjacents est effectué par des arêtes entre les sommets de ces derniers.

Dans l'algorithme des k-lignes, la droite  $s_i$  est définie comme  $\mathbf{s}_i = \{\mathbf{s}_i(t) | t \in \mathcal{R}\}$ , où  $\mathbf{s}_i(t) = \mathbf{c}_i + p_i t$ . La distance entre un point et la droite est donnée par :

$$d(\mathbf{x}, \mathbf{s}_i) = \inf_{t \in \mathcal{R}} \|\mathbf{s}_i(t) - \mathbf{x}\| \quad (4.14)$$

Soit  $X$  l'ensemble des échantillons avec  $N$  mesures et  $m$  variables. Des régions  $R_1, \dots, R_r$  sont définies par :

$$R_i = \left\{ \mathbf{x} \in X \mid i = \arg \min_j d(\mathbf{x}, \mathbf{s}_j) \right\} \quad (4.15)$$

Dans l'algorithme k-lines, l'objectif est de trouver  $r$  droites  $s_1, \dots, s_r$  qui minimisent :

$$\sum_{i=1}^r \sum_{\mathbf{x} \in R_i} d(\mathbf{x}, \mathbf{s}_i)^2 \quad (4.16)$$

Ainsi, pour déterminer les droites qui sont des optima locaux de (4.16), Verbeek [101] propose d'exploiter l'algorithme des k-lignes. Commencer avec des centres et des directions aléatoires des  $k$  droites, et répéter les deux étapes suivantes jusqu'à convergence :

1. déterminer les régions  $R_i$
2. déterminer la droite pour chaque région  $R_i$  comme étant la première composantes principale de la matrice de corrélation des données  $\in R_i$ .

Puisque on cherche à construire une courbe polygonale, il ne faut pas chercher des droites mais des segments de droite. Ainsi, on doit remplacer l'étape 2 de l'algorithme. Au lieu d'utiliser la première composante principale, on doit utiliser un segment de cette dernière. Ainsi, on utilise le plus petit segment de la première composantes telle que les projections des points de la région correspondante sur cette composante soient incluses dans le segment.

Cependant, il reste à résoudre le problème de liaison entre les différents segments de droite obtenus pour obtenir une courbe polygonale. Pour cela, Vereek [101] définit un graphe totalement connecté  $G = (R, E)$ , où l'ensemble de sommets est constitué de  $2k$  points des k-segments. Ainsi, deux sommets adjacents sont connecté par une arête minimisant la distance euclidienne entre ces sommets tout en minimisant la distance entre la courbe obtenues et l'ensemble des points [101].

Il faut noter qu'une fois la courbe principale est obtenue, chaque point de la composante principale non-linéaire correspondante est calculé, comme dans le cas de l'algorithme de Hastie, comme étant la longueur d'arc le long de la courbe.

Les figures (fig 4.5), (fig 4.6), (fig 4.7) et (fig 4.8) présentent, respectivement, les courbes obtenues en utilisant 1, 2, 3 et 4 segments de droite. La figure (fig 4.9) présente la courbe représentant les relations non linéaires entre les trois variables  $x_1, x_2, x_3$  ainsi que l'estimation de cette courbe en utilisant l'algorithme des courbes principales développé par Verbeek [101]. Pour l'estimation de cette courbe nous avons utilisé cinq segments de droite.

L'approche des courbes principales a été présentée comme une technique de généralisation de l'ACP linéaire monodimensionnelle. Une méthode des surfaces principales (cas de dimension 2) a été également présentée par Hastie [38]. Cependant, pour des problèmes de dimension plus élevée, la paramétrisation des hyper-surfaces est plus complexe. L'extension de l'algorithme des courbes

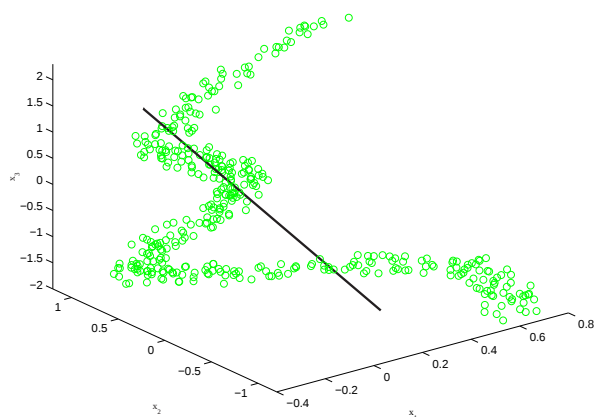


Figure 4.5 – Utilisation d'un seul segment de droite

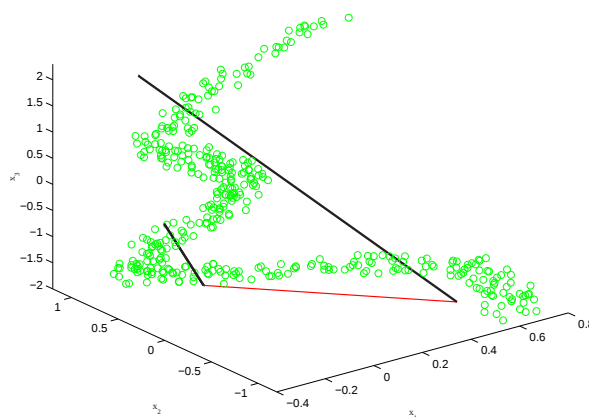


Figure 4.6 – Utilisation de deux segments de droite

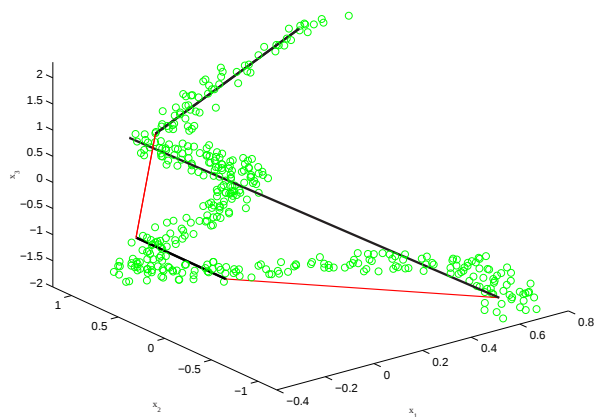


Figure 4.7 – Utilisation de trois segments de droite

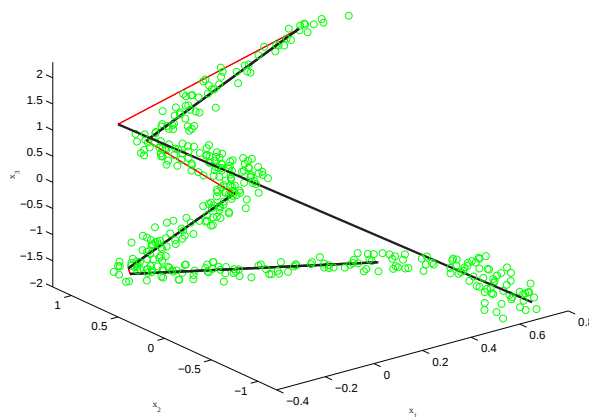


Figure 4.8 – Utilisation de quatre segments de droite

principales dans ce cas a été proposée [60] ; le nouvel algorithme est basé sur une extraction des courbes principales d'une manière séquentielle. Ainsi, l'équation (4.2) peut s'écrire :

$$X = \mathcal{F}_1(t_1) + \mathcal{F}_2(t_2) + \dots + \mathcal{F}_\ell(t_\ell) + E \quad (4.17)$$

Les composantes principales non linéaires sont calculées d'une manière séquentielle et à chaque itération une composante non linéaire est calculée. La recherche d'une composante est équivalente à la recherche d'une courbe principale en minimisant la distance entre cette courbe et la projection des données sur cette courbe. Ensuite l'approximation par la courbe des variables est retirée des variables originelles pour itérer l'algorithme. Ainsi, pour la recherche de  $\ell$  composantes non-linéaire on utilise la procédure suivante :

1. poser  $i=1$ ,
2. rechercher la  $i^{eme}$  courbe principale  $X = \mathcal{F}_i(t_i) + E_i$ ,
3. calculer l'erreur d'estimation  $E_i = X - \mathcal{F}(t_i)$ ,
4. si  $i = \ell$  arrêter, sinon poser  $X = E_i$  et  $i = i + 1$  et aller à l'étape 2.

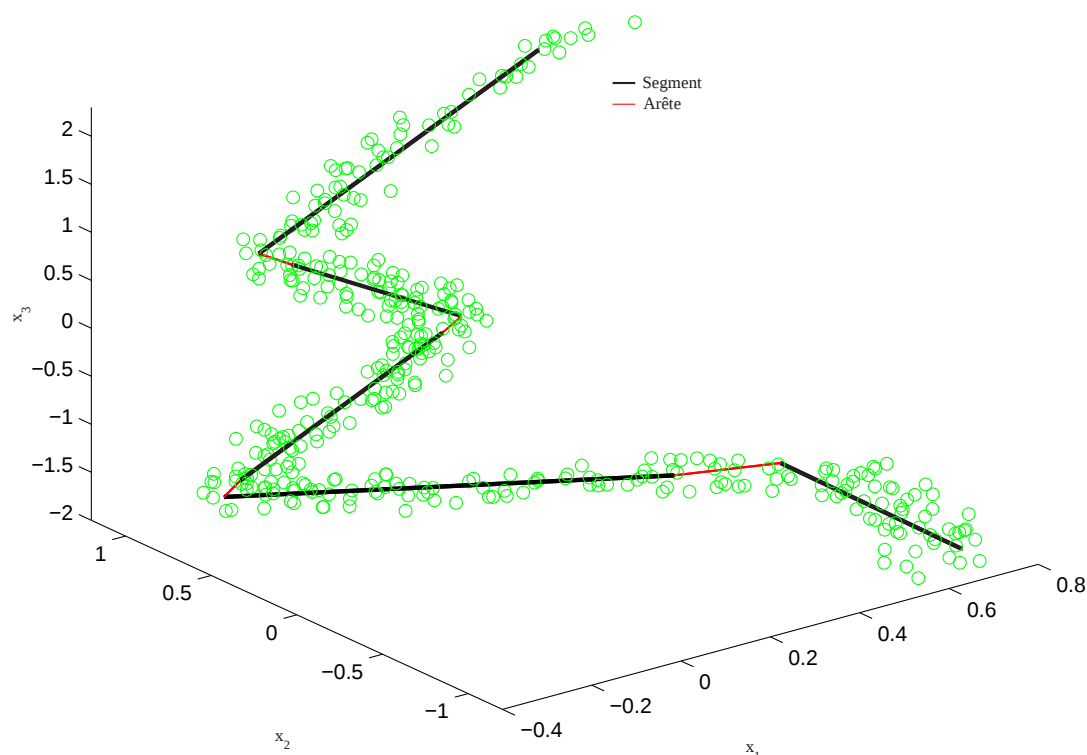


Figure 4.9 – Courbe principale pour l'exemple présenté avec l'algorithme de Verbeek en utilisant cinq segments de droite

Cette procédure peut s'appliquer à l'algorithme de Verbeek pour l'extraction de plusieurs composantes principales non-linéaires.

L'inconvénient majeur des courbes principales est qu'elles ne donnent pas un modèle de représentation des composantes principales non linéaires. A chaque point de l'ensemble des mesures lui est associée une valeur de  $t$  correspondant à la composante non linéaire pour ce point. Ainsi, pour de nouvelles observations, il est très difficile voire impossible de calculer les valeurs des composantes non linéaires correspondantes.

Cependant, pour le diagnostic, il est préférable d'avoir une représentation des composantes principales non linéaires permettant le calcul de la valeur des composantes non linéaires pour de nouvelles observations. Les réseaux de neurones artificiels sont reconnus pour leur efficacité à identifier les relations non linéaires entre variables. Dans la section suivante, seront exposés les algorithmes d'ACPNL à base de réseaux de neurones.

## 4.4 Approches neuronales de l'ACPNL

L'analyse en composantes non linéaires (ACPNL) à base de réseaux de neurones, a connu ces dernières années un regain d'intérêt considérable et a été largement utilisée dans le domaine du diagnostic [15, 94, 13]. Dans cette partie, nous allons présenter les différentes architectures de réseaux pour l'extraction des composantes principales non linéaires.

#### 4.4.1 ACPNL par réseau à cinq couches

L'analyse en composantes principales non linéaires (ACP NL) a eu un intérêt particulier ces dernières années [64, 104, 10, 102, 105, 67, 42, 54]. La plupart des chercheurs utilise une approche neuronale pour le calcul du modèle ACPNL qui a été proposée par Kramer [56].

Kramer [56] a proposé une ACPNL comme une extension de l'ACP linéaire. Dans le cas d'une seule composante principale non-linéaire, l'architecture d'un tel réseau est illustrée sur la figure (fig 4.10).

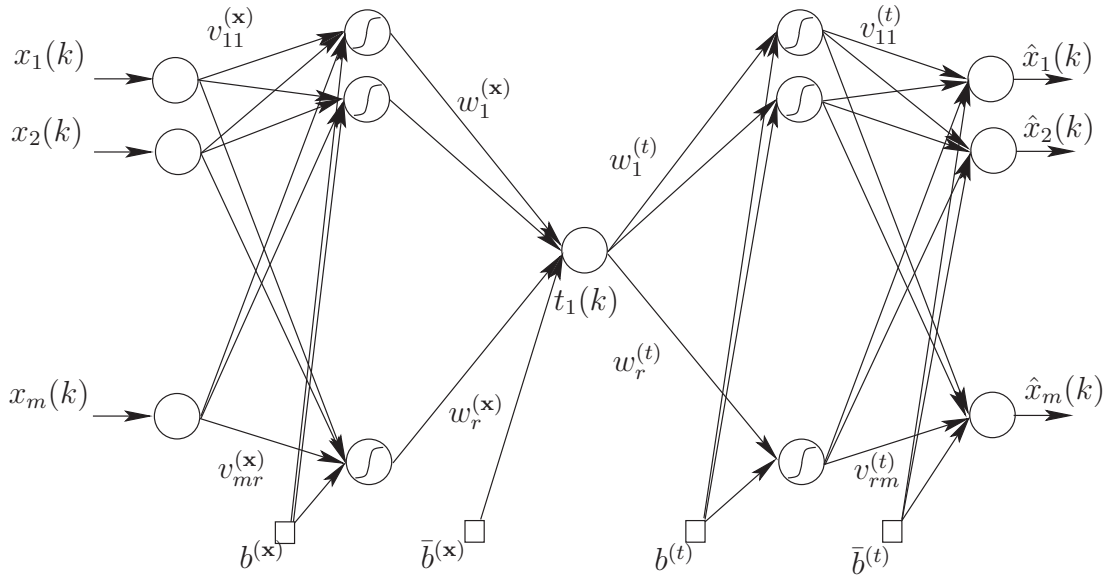


Figure 4.10 – Réseau à cinq couches pour l'extraction d'une seule composante principale non linéaire

Pour effectuer l'ACP NL, le réseau de neurones (fig 4.10) contient trois couches entre les variables d'entrées et de sorties. Une fonction de transfert  $\mathcal{G}_1$  réalise une projection du vecteur  $\mathbf{x}$ , vecteur d'entrée de dimension  $m$ , vers la première couche cachée (couche de codage), représentée par  $h_j^{(\mathbf{x})}$  ( $j = 1, \dots, r$ ), un vecteur colonne de dimension  $r$  et  $r$  représente le nombre de neurones dans la première couche cachée :

$$h_j^{(\mathbf{x})} = \mathcal{G}_1 \left( \left( V^{(\mathbf{x})} \mathbf{x} + b^{(\mathbf{x})} \right) \right) = \mathcal{G}_1 \left( \sum_{i=1}^m v_{ij}^{(\mathbf{x})} x_i + b_j^{(\mathbf{x})} \right) \quad (4.18)$$

$V^{(\mathbf{x})}$  est la matrice des poids de dimension  $(r \times m)$ ,  $b^{(\mathbf{x})}$  un vecteur contenant les  $r$  paramètres de biais. La deuxième fonction de transfert  $\mathcal{G}_2$  projette les données de la première couche cachée vers la couche d'étranglement "bottleneck layer" contenant un seul neurone, qui représente la composante principale non linéaire  $t$ . La fonction de transfert  $\mathcal{G}_1$  est généralement non linéaire (on utilise la fonction tangente hyperbolique ou la fonction sigmoïde), tandis que la fonction  $\mathcal{G}_2$  représente la fonction identité ( $\mathcal{G}_2(\mathbf{x}) = \mathbf{x}$ ) :

$$t = \mathcal{G}_2 \left( \mathbf{w}^{(\mathbf{x})} h^{(\mathbf{x})} + \bar{b}^{(\mathbf{x})} \right) = \sum_{j=1}^r w_j^{(\mathbf{x})} h_j^{(\mathbf{x})} + \bar{b}^{(\mathbf{x})} \quad (4.19)$$

Ensuite, la fonction de transfert  $\mathcal{G}_3$ , qui est une fonction non linéaire, projette les données à

partir de  $t$  vers la dernière couche cachée (couche de décodage)  $h_j^{(t)}$  ( $j = 1, \dots, r$ ) et  $r$  représente le nombre de neurones dans la troisième couche cachée :

$$h_j^{(t)} = \mathcal{G}_3 \left( \left( \mathbf{w}^{(t)} t + b^{(t)} \right)_j \right) = \mathcal{G}_3 \left( w_j^{(t)} t + b_j^{(t)} \right) \quad (4.20)$$

la dernière fonction de transfert  $\mathcal{G}_4$  est la fonction identité et projette les données à partir de  $h^{(t)}$  vers  $\hat{\mathbf{x}}$ , le vecteur de sortie de dimension  $m$  :

$$\hat{x}_i = \mathcal{G}_4 \left( \left( \mathbf{V}^{(t)} h^{(t)} + \bar{b}^{(t)} \right)_i \right) = \sum_{j=1}^r v_{ji}^{(t)} h_j^{(t)} + \bar{b}_i^{(t)} \quad (4.21)$$

La fonction coût  $E = \|\mathbf{x} - \hat{\mathbf{x}}\|^2$  est minimisée pour trouver les valeurs optimales de  $V^{(\mathbf{x})}$ ,  $b^{(\mathbf{x})}$ ,  $\bar{b}^{(\mathbf{x})}$ ,  $\mathbf{w}^{(t)}$ ,  $b^{(t)}$ ,  $V^{(t)}$  et  $\bar{b}^{(t)}$ .

Une fois la structure générale définie, il reste à déterminer la structure précise du modèle neuronal (architecture du réseau). Pour cela, il faut déterminer le nombre de couches cachées nécessaire et le nombre de neurones dans chaque couche cachée. Pour la détermination du nombre de couches cachées, Funahashi [21] et Cybenko [8] ont montré que toute fonction continue peut être approximée par un réseau de neurones à trois couches utilisant une fonction d'activation sigmoïdale pour les neurones de la couche cachée et une fonction d'activation linéaire pour les neurones de la couche de sortie. Le nombre de neurones dans la couche cachée est généralement déterminé en effectuant une validation croisée sur un jeu de validation. Il existe des algorithmes permettant de construire itérativement la couche cachée, le contrôle de la croissance du réseau est effectué par une validation croisée. D'autres méthodes utilisent exactement un chemin inverse des précédentes, partant d'un réseau avec un grand nombre de neurones dans la couche cachée ; les connexions jugées inutiles sont éliminées progressivement jusqu'à l'obtention d'une structure satisfaisante. Une synthèse des techniques d'élimination est présentée dans le papier de Kerling [53]. Cependant, il faut noter ici que l'on ne s'intéresse pas à ce problème.

#### 4.4.1.1 Apprentissage du réseau

La plupart des algorithmes d'apprentissage des réseaux de neurones sont des algorithmes d'optimisation ; ils cherchent à minimiser, par des méthodes d'optimisation non linéaires, une fonction de coût. Cette optimisation se fait de manière itérative, en modifiant les poids en fonction du gradient de la fonction de coût (qui, dans le cas d'un apprentissage supervisé, constitue une mesure de l'écart entre les réponses réelles du réseau et ses réponses désirées). Les poids sont initialisés aléatoirement avant l'apprentissage, puis modifiés itérativement jusqu'à obtention d'un compromis satisfaisant entre la précision de l'approximation sur l'ensemble d'apprentissage et la précision de l'approximation sur l'ensemble de validation.

On note  $E(k)$  l'erreur quadratique obtenue à partir des erreurs obtenues sur les  $m$  neurones de sortie en présence de l'exemple  $k$  :

$$E(k) = \frac{1}{2} \sum_{i=1}^m (x_i(k) - \hat{x}_i(k))^2 \quad (4.22)$$

La minimisation de  $E(k)$  par rapport aux poids  $v_{ji}^{(t)}$  est donnée par :

$$\frac{\partial E(k)}{\partial v_{ji}^{(t)}} = \frac{\partial E(k)}{\partial \hat{x}_i(k)} \frac{\partial \hat{x}_i(k)}{\partial v_{ji}^{(t)}} \quad (4.23)$$

$$\frac{\partial E(k)}{\partial \hat{x}_i(k)} = -(x_i(k) - \hat{x}_i(k)) \text{ et } \frac{\partial \hat{x}_i(k)}{\partial v_{ji}^{(t)}} = h_j^{(t)}.$$

Le gradient de  $E(k)$  calculé sur la couche de sortie est donné par :

$$\frac{\partial E(k)}{\partial v_{ji}^{(t)}} = -(x_i(k) - \hat{x}_i(k)) h_j^{(t)} \quad (4.24)$$

Ainsi, les modifications des poids sont données par :

$$\Delta v_{ji}^{(t)} = \mu (x_i(k) - \hat{x}_i(k)) h_j^{(t)} \quad (4.25)$$

De même pour les couches intermédiaires, commençons par calculer le gradient de  $E(k)$  par rapport à  $w_j^{(t)}$  :

$$\frac{\partial E(k)}{\partial w_j^{(t)}} = \frac{\partial E(k)}{\partial \hat{x}_i(k)} \frac{\partial \hat{x}_i(k)}{\partial h_j^{(t)}} \frac{\partial h_j^{(t)}}{\partial w_j^{(t)}} \quad (4.26)$$

$$\frac{\partial E(k)}{\partial w_j^{(t)}} = -\mathcal{G}'_3 \left( w_j^{(t)} t + b_j^{(t)} \right) t \sum_{i=1}^m (x_i(k) - \hat{x}_i(k)) v_{ji}^{(t)} \quad (4.27)$$

Les modifications de  $w_j^{(t)}$  sont données par :

$$\Delta w_j^{(t)} = \mu \mathcal{G}'_3 \left( w_j^{(t)} t + b_j^{(t)} \right) t \sum_{i=1}^m (x_i(k) - \hat{x}_i(k)) v_{ji}^{(t)} \quad (4.28)$$

Les mêmes opérations de dérivation et de calcul des modifications seront appliquées pour les poids  $w_j^{(\mathbf{x})}$  et  $v_{ij}^{(\mathbf{x})}$ , ce qui nous permet d'écrire :

$$\frac{\partial E(k)}{\partial w_j^{(\mathbf{x})}} = \frac{\partial E(k)}{\partial t} \frac{\partial t}{\partial w_j^{(\mathbf{x})}} \quad (4.29)$$

La quantité  $\frac{\partial t}{\partial w_j^{(\mathbf{x})}}$  peut être exprimée comme  $\frac{\partial t}{\partial w_j^{(\mathbf{x})}} = h_j^{(\mathbf{x})}$  et ainsi on obtient l'expression du gradient de  $E(k)$  par rapport à  $w_j^{(\mathbf{x})}$  :

$$\frac{\partial E(k)}{\partial w_j^{(\mathbf{x})}} = -h_j^{(\mathbf{x})} \sum_{i=1}^r \mathcal{G}'_3 \left( w_j^{(t)} t + b_j^{(t)} \right) w_j^{(t)} \sum_{i=1}^m (x_i(k) - \hat{x}_i(k)) v_{ji}^{(t)} \quad (4.30)$$

Finalement, le gradient de  $E(k)$  par rapport aux poids  $v_{ij}^{(\mathbf{x})}$  est donné par :

$$\frac{\partial E(k)}{\partial v_{ij}^{(\mathbf{x})}} = \frac{\partial E(k)}{\partial t} \frac{\partial t}{\partial h_j^{(\mathbf{x})}} \frac{\partial h_j^{(\mathbf{x})}}{\partial v_{ij}^{(\mathbf{x})}} \quad (4.31)$$

$$\frac{\partial E(k)}{\partial v_{ij}^{(\mathbf{x})}} = -\mathcal{G}'_1 \left( \sum_{i=1}^m v_{ij}^{(\mathbf{x})} x_i + b_j^{(\mathbf{x})} \right) x_i \sum_{p=1}^r w_p^{(\mathbf{x})} \sum_{j=1}^r \mathcal{G}'_3 \left( w_j^{(t)} t + b_j^{(t)} \right) w_j^{(t)} \sum_{i=1}^m (x_i(k) - \hat{x}_i(k)) v_{ji}^{(t)}$$

Ainsi, on a obtenu les gradients de  $E(k)$  par rapport aux différents poids du réseau et la règle de modification de ces poids est donnée par  $\Delta w = -\mu \frac{\partial E}{\partial w}$  où  $\mu$  est le pas d'adaptation.

Malthose [64] propose pour l'apprentissage du réseau de neurones à cinq couches pour produire ces entrées de minimiser la fonction de coût suivante :

$$\min_{\mathcal{F}, \mathcal{G}} \sum_{k=1}^N \left( \|\mathbf{x}(k) - \mathcal{F}(\mathcal{G}(\mathbf{x}(k)))\|^2 + \left( \sum_{i=1}^m \left( \frac{\partial \mathcal{F}(\mathbf{x}(k))}{\partial t} \right)^2 - 1 \right) \right) \quad (4.32)$$

Le terme additif force le réseau à produire une courbe avec la propriété (4.10).

L'approximation de la courbe de l'exemple traité avec un réseau à cinq couches, avec six neurones dans chaque couche cachée, est illustrée sur la figure (fig 4.11). Ainsi, on obtient une estimation permettant d'expliquer environ 97% des corrélations totales des variables.

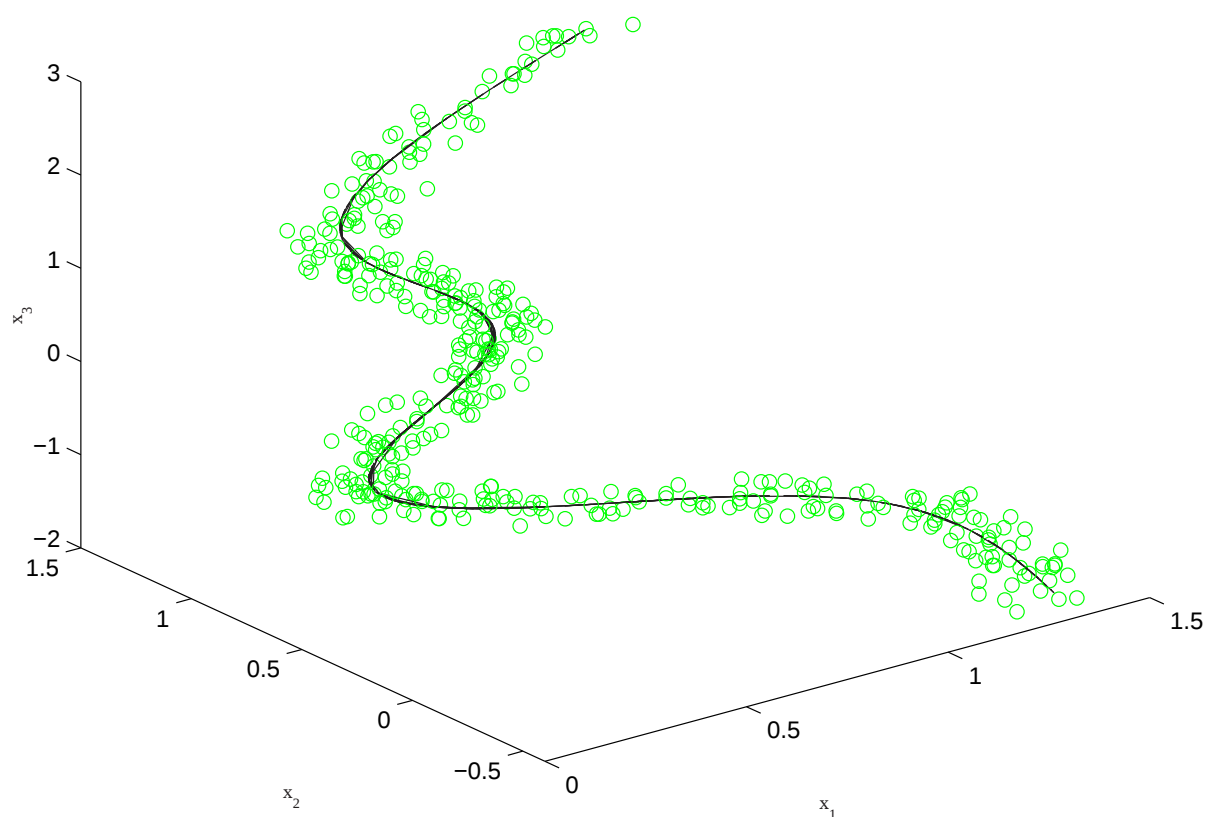


Figure 4.11 – Mesures et estimation avec la première composante non linéaire avec un réseau à cinq couches

Il faut noter que l'extraction des composantes principales peut se faire de deux façons.

La première consiste à extraire les composantes séquentiellement en ayant qu'un seul neurone dans la couche du milieu "bottleneck layer" (ACPNL séquentielle) (fig 4.12). La seconde consiste à extraire les  $\ell$  composantes désirées simultanément en insérant  $\ell$  neurones dans la couche du milieu (ACPNL parallèle ou simultanée) (fig 4.13).

L'ACPNL parallèle nécessite, avant l'apprentissage du réseau, la détermination du nombre de composantes non linéaires à retenir  $\ell$ . Tandis que dans le cas de l'ACPNL séquentielle le réseau est entraîné initialement avec une seule composante principale non linéaire. Après estimation des données à partir de cette première composante non linéaire, on doit soustraire le résultat obtenu

de l'ensemble des données de départ et l'opération d'extraction d'une deuxième composante non linéaire est effectuée sur les résidus obtenus. Cette procédure peut être répétée jusqu'à ce que le nombre de composantes voulu soit atteint ou l'erreur d'estimation inférieure à un certain seuil choisi a priori.

#### Algorithme séquentiel (fig 4.12)

1. Projeter les données  $X$  dans un espace de dimension 1 des composantes principales en utilisant une structure de réseau à cinq couches.
2. Apprendre le réseau et estimer  $\hat{X}$  par projection inverse.
3. Calculer l'erreur d'estimation  $E = X - \hat{X}$ . Répéter les étapes 1 et 2 en prenant la matrice  $E$  à la place de  $X$  jusqu'à ce qu'un pourcentage de la variance de  $X$  soit capturé.

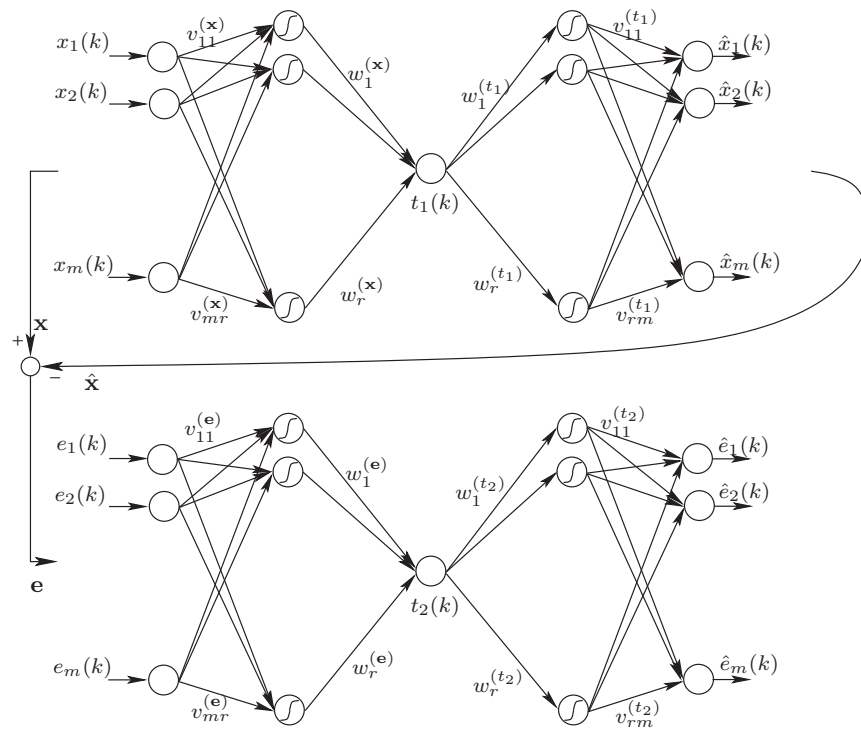
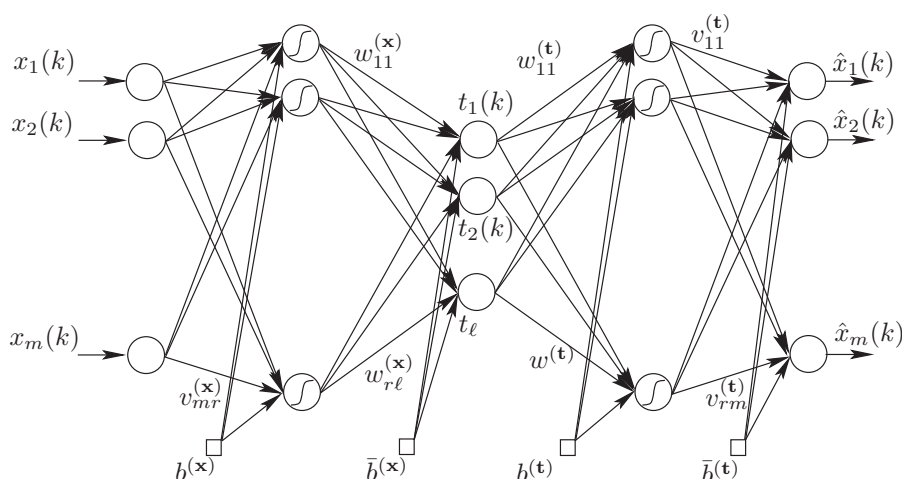


Figure 4.12 – Extraction séquentielle des composantes principales non linéaires

#### Algorithme parallèle (fig 4.13)

1. Projeter les données  $X$  dans un espace des composantes principales de dimension  $\ell$  en utilisant une structure de réseau à cinq couches.
2. Apprendre le réseau et estimer les données  $\hat{X}$  par projection inverse.

Les deux algorithmes ont leurs avantages et inconvénients. L'avantage principal de l'algorithme parallèle est qu'il est plus rapide à apprendre en le comparant à l'algorithme séquentiel : pour  $\ell$  composantes, un seul réseau de neurones est nécessaire dans le cas parallèle, alors qu'il faut  $\ell$  réseaux dans le cas séquentiel. L'inconvénient de l'algorithme parallèle réside dans la première étape. Cette étape suppose que le nombre des composantes principales doit être connu a priori.

Figure 4.13 – Extraction simultanée des  $\ell$  composantes principales non linéaires

Cependant, en plus du nombre de paramètres à calculer il y a le problème de convergence de l'apprentissage de ce type de réseaux. Ainsi, pour simplifier le problème d'apprentissage d'un tel réseau, une solution consiste à le diviser en deux sous réseaux de trois couches et d'apprendre chaque réseau séparément. Le premier problème rencontré avec cette solution est que l'on ne connaît pas les sorties du premier réseau (qui représentent normalement les entrées du second réseau).

Plusieurs approches ont été proposées pour résoudre ce problème que nous allons présenter dans la suite.

#### 4.4.2 ACPNL par optimisation des entrées du réseau (Input Training network)

Tan et Mavrovouniotis [92] proposent une approche pour l'analyse en composantes principales non linéaires (ACPNL) basée sur le concept d'apprentissage des entrées (représentant les composantes principales recherchées) d'un réseau de neurones IT-net (Input Training Network). L'architecture d'un réseau IT-net est illustrée sur la figure (fig 4.14). C'est un réseau à trois couches dont une couche cachée. La couche de sortie est composée de  $m$  neurones correspondant à la dimension des données  $\mathbf{x}$ . La couche d'entrée contient  $\ell$  neurones correspondant au nombre de composantes principales non linéaires  $\mathbf{t}$ .

Au lieu d'entraîner un réseau auto-associatif à cinq couches, il est préférable d'entraîner seulement une partie de ce réseau composée de trois couches (sous-réseau de décompression). Entraîner un tel sous-réseau est intéressant et peut être effectué par extension de l'algorithme de rétro-propagation, puisque la fonction d'erreur est bien définie. La différence entre l'apprentissage de ce réseau et un réseau multi-couches ordinaire est que les entrées de ce réseau ne sont pas connues, car elles représentent les composantes principales recherchées. Donc, dans la phase d'apprentissage il faut ajuster non seulement les paramètres internes du réseau mais également les valeurs des entrées par minimisation de l'erreur de sortie du réseau.

Il faut noter que Tan [92] a montré que l'apprentissage de ce réseau avec des neurones linéaires et sans couche cachée est équivalent à la méthode de puissance (Power method) [86] utilisée pour le calcul des vecteurs propres de la matrice de covariance des données. Il a montré également que l'apprentissage d'un réseau IT avec une seule entrée et avec une seule couche cachée non linéaire

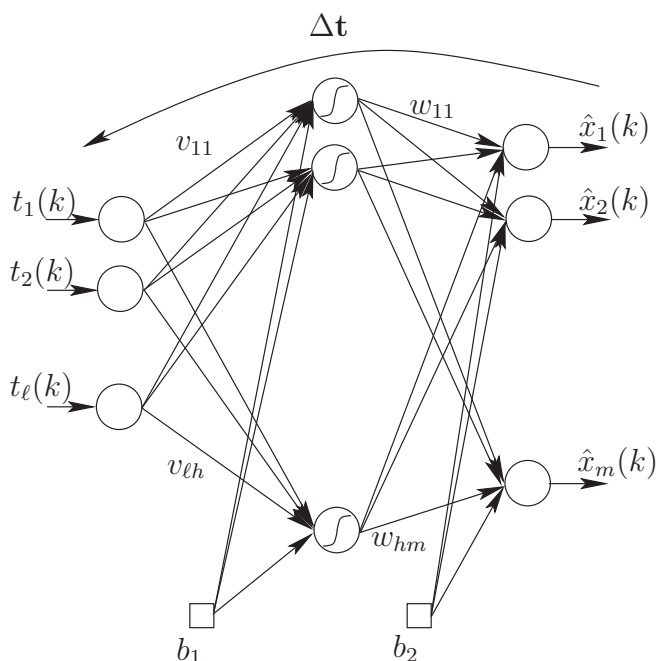


Figure 4.14 – Réseau avec apprentissage des entrées (IT)

est équivalent à l'approche des courbes principales de Hastie [39].

L'apprentissage d'un réseau IT permet de déterminer la fonction de projection entre les composantes principales non linéaires et les données du processus. Toutefois, le modèle permettant de calculer, à partir des nouvelles observations, les composantes principales non linéaires correspondantes n'est pas défini. Pour résoudre ce problème, Tan et Mavrouniotis [92] proposent un algorithme d'optimisation en ligne des entrées du réseau IT pour le calcul des composantes principales non linéaires pour les nouvelles observations. La limitation de cette approche est qu'une procédure d'optimisation non linéaire est nécessaire pour la recherche des composantes principales non linéaires permettant de donner une meilleure approximation des nouvelles observations.

#### 4.4.3 Réseaux à trois couches et courbes principales

En combinant la méthodologie des courbes principales [39] et les réseaux de neurones auto-associatifs, Dong et McAvoy [12] proposent une procédure pour l'identification du modèle ACP non linéaire. En considérant le réseau auto-associatif comme deux sous-réseaux avec une seule couche cachée chacun, représentant le réseau de compression et de décompression respectivement, Dong et McAvoy [12] proposent d'apprendre les deux sous-réseaux séparément. La méthode proposée nécessite trois étapes :

- identification des composantes non linéaires par application successive de l'algorithme des courbes principales de Hastie [39] sur les données du processus,
- apprentissage du réseau à trois couches ayant comme entrées les données du processus et les composantes principales comme sorties,
- apprentissage du second réseau à trois couches pour l'estimation des données originelles, ayant les composantes principales non linéaires comme entrées.

L'apprentissage des deux réseaux à trois couches peut être effectué par n'importe quel algorithme adéquat : gradient conjugué ou Levenberg-Marquardt [61, 66].

L'illustration de cette approche est effectuée en utilisant notre exemple en utilisant l'algorithme de Verbeek [101] pour le calcul de la composante principale utilisée pour l'apprentissage des deux sous-réseaux. Les figures (fig 4.22) et (fig 4.23) présentent l'estimation de la courbe de l'exemple traité en combinant deux sous-réseaux de neurones avec six neurones dans chaque couche cachée. L'estimation obtenue explique environ 97% des corrélations des variables que se soit sur le jeu de données d'identification ou sur le jeu de validation.

De plus, à partir des deux figures (fig 4.17) et (fig 4.18) représentant, respectivement, la première composante principale non-linéaire obtenue en utilisant le réseau à cinq couches et l'approche des courbes principales, il est clair que les deux composantes sont quasiment identiques d'où l'intérêt de cette approche de combinaison. Pour un cas plus compliqué, on peut extraire plusieurs composantes on utilisant l'algorithme de Verbeek.

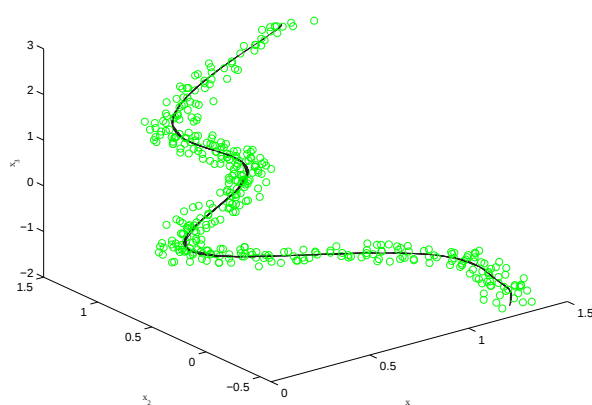


Figure 4.15 – Mesures et estimation de la courbe de l'exemple 1 en combinant les courbes principales et deux réseaux MLP a trois couches (jeu d'identification)

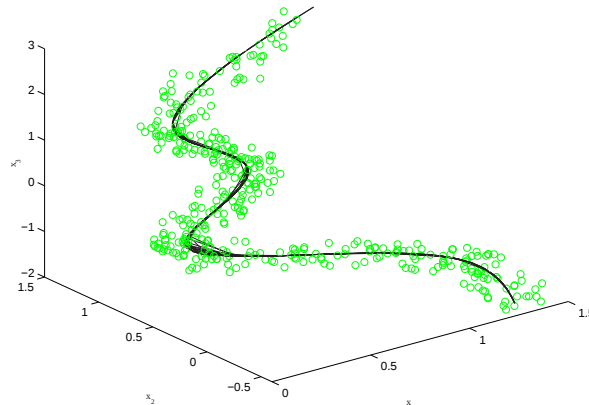


Figure 4.16 – Mesures et estimation de la courbe de l'exemple 1 en combinant les courbes principales et deux réseaux MLP a trois couches (jeu de validation)

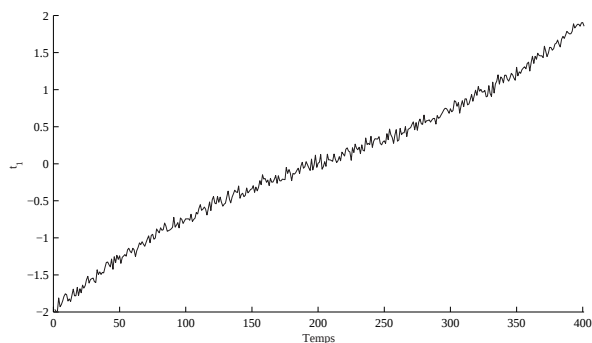


Figure 4.17 – Evolution de la première composante principale non-linéaire obtenue avec le réseau de neurones à cinq couches

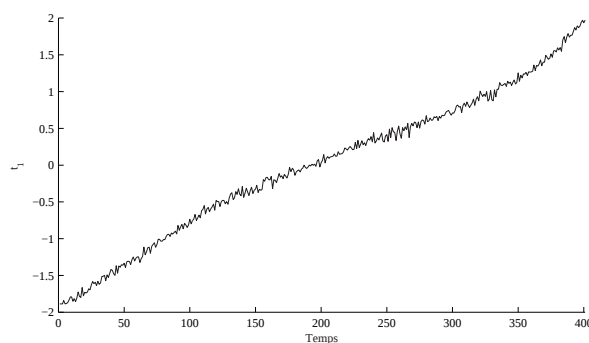


Figure 4.18 – Evolution de la première composante principale non-linéaire obtenue avec l'algorithme de Verbeek

Il est clair que le problème d'apprentissage est devenu plus simple en le comparant avec l'apprentissage du réseau à cinq couches ou le réseau IT. Toutefois, on peut encore simplifier l'apprentissage en utilisant les réseaux RBF.

#### 4.4.4 Réseaux de fonctions à base radiale (RBF)

Un grand avantage des réseaux RBF, par rapport aux autres types de réseaux, est que de bonnes performances sont obtenues par un apprentissage qui est généralement beaucoup plus rapide que l'apprentissage des autres modèles connexionnistes pour des problèmes de régression. Il existe plusieurs méthodes dont le principe est d'approximer un comportement désiré par une collection de fonctions, appelées fonctions noyau (kernel functions). La méthode RBF a comme particularité que ces fonctions-noyau sont locales, c'est-à-dire qu'elles ne donnent des réponses utiles que pour un domaine de valeurs restreint. Ce domaine est défini autour d'un point, le centre (ou noyau). En général, chaque fonction-noyau est décrite par deux paramètres : la position de son centre et la taille du domaine. Pour approximer un comportement donné, les fonctions-noyau sont assemblées pour couvrir de leurs domaines l'ensemble des données d'entrée. Ces fonctions sont ensuite pondérées et leurs valeurs sont sommées pour produire une valeur de sortie. Les fonctions-noyau peuvent être définies, de manière générale, par :

$$\phi(\mathbf{x}, c) = \phi(\|\mathbf{x} - \mathbf{c}\|; \Sigma) \quad (4.33)$$

où  $\phi(\cdot)$  est une fonction-noyau,  $\|\cdot\|$  est une norme (en général la norme Euclidienne ou la distance de Mahalanobis) et  $\mathbf{c}$  est le centre de la fonction noyau.

Pour définir le modèle ACPNL à base de réseau RBF, considérons les deux réseaux RBF donnés par les deux figures (fig 4.19) et (fig 4.20). Le premier réseau définit la projection non-linéaire pour le calcul des composantes principales non-linéaires  $t$  à partir des données  $\mathbf{x}$ . Ainsi, la sortie  $t$  du réseau peut s'écrire sous la forme :

$$t = \mathcal{G}(\mathbf{x}) \quad (4.34)$$

$$t = \sum_{i=1}^r w_i \phi_i(\mathbf{x}) = \mathbf{w}^T \Phi(\mathbf{x}) \quad (4.35)$$

où  $\mathbf{w} \in \mathbb{R}^r$  est le vecteur des poids de la couche de sortie du premier réseau,  $\Phi \in \mathbb{R}^r$  est le vecteur des variables transformées et  $r$  représente le nombre de transformations ou le nombre de neurones dans la couche cachée.

Le second réseau définit la transformation inverse permettant de calculer l'estimation  $\hat{\mathbf{x}}$  à partir de  $t$ . Ainsi, on a :

$$\hat{\mathbf{x}} = \mathcal{F}(t) \quad (4.36)$$

$$\hat{\mathbf{x}} = \sum_{j=1}^r v_j \psi_j(t) + v_0 = \mathbf{V}^T \Psi(t) + v_0 \quad (4.37)$$

où  $\psi_j$  ( $j = 1, \dots, r$ ) représente les fonctions noyau et  $V^T = [v_1 \dots v_r] \in \mathbb{R}^{m \times r}$  est la matrice des poids.  $v_0$  est un terme de biais et  $r$  représente le nombre de noyaux.

Dans la suite, nous allons considérer deux cas d'apprentissage du réseau RBF pour le calcul du modèle ACPNL en utilisant deux sous-réseaux RBF à trois couches. Les deux sous-réseaux sont représentés sur les figures (fig 4.19) et (fig 4.20).

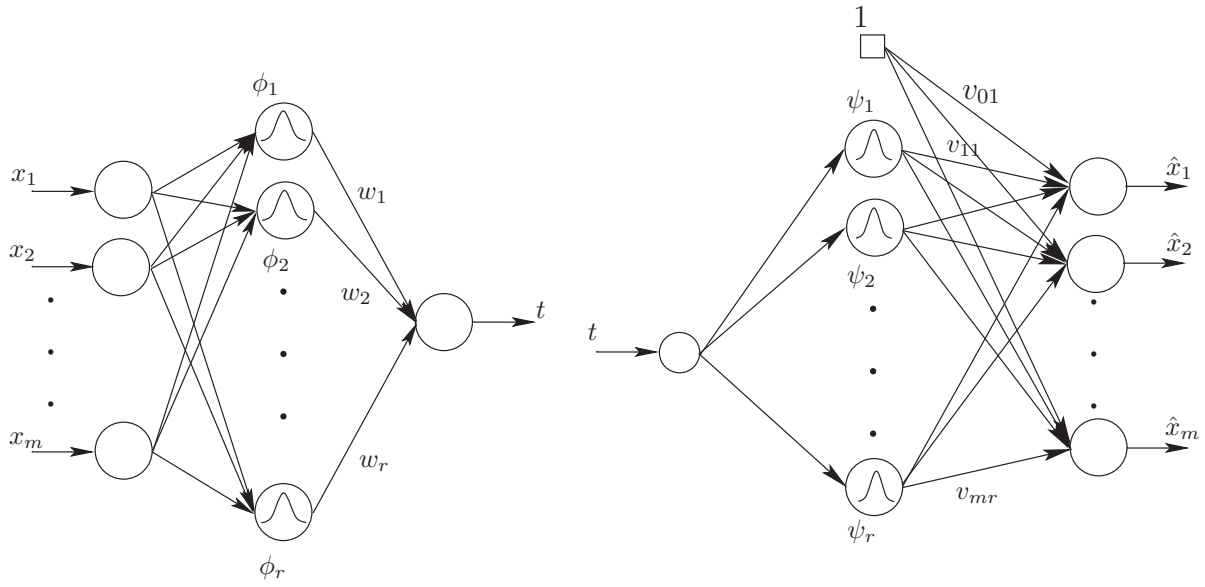


Figure 4.19 – Réseau RBF pour la compression des données (projection)

Figure 4.20 – Réseau RBF pour la décompression des données (projection inverse)

Puisque les composantes principales non-linéaires  $t$ , représentant à la fois les sorties du premier réseau (fig 4.19) et les entrées du second réseau (fig 4.20), ne sont pas connues. Webb [104] propose de calculer les composantes principales à partir du premier réseau en maximisant la variance de sortie de ce dernier.

L'approche que nous proposons, utilise l'algorithme de Verbeek [101] pour le calcul des composantes principales non linéaires, une fois les composantes connues l'apprentissage des deux réseaux se ramène à un problème de régression linéaire par rapport aux poids.

#### 4.4.4.1 Réseau RBF avec maximisation de la variance de $t$

Considérons le réseau à fonctions de base radiale illustré par la figure (fig 4.19). Pour simplifier en présentera le cas d'une seule composantes principale. Il faut noter ici que l'on cherche à trouver les poids  $w_i$ , ( $i = 1, \dots, r$ ) en supposant pour l'instant que les fonctions  $\phi_i$  sont identifiées (centres et dispersions optimisées).

Ainsi, d'après l'expression (4.35), On cherche  $\mathbf{w}$  qui maximise la variance de  $t$  donnée par :

$$Var(t) = \frac{1}{N} \sum_{k=1}^N (t(k) - \bar{t}) (t(k) - \bar{t})^T = \mathbf{w}^T \mathbf{A} \mathbf{w} \quad (4.38)$$

où  $\bar{t}$  désigne la moyenne de  $t$  sur l'ensemble d'apprentissage et  $\mathbf{A}$  est la matrice de covariance des variables  $\phi(\mathbf{x})$  qui est donnée par :

$$\mathbf{A} = \frac{1}{N} \sum_{k=1}^N ((\Phi_k - \bar{\Phi}) (\Phi_k - \bar{\Phi})^T) \quad (4.39)$$

où  $\Phi_k = (\phi_1(\mathbf{x}(k)), \dots, \phi_r(\mathbf{x}(k)))^T$  et  $\bar{\Phi}$  est la moyenne de  $\Phi$ .

Imposons la contrainte de normalisation sur  $\mathbf{w}$  :

$$\mathbf{w}^T \mathbf{B} \mathbf{w} = 1 \quad (4.40)$$

où  $\mathbf{B}$  est une matrice symétrique.

Cette contrainte impose une condition sur le gradient de  $t$  [104, 111]. La condition imposée est que l'amplitude quadratique moyenne du gradient de  $t$  soit égale à 1. Dans le cas linéaire ( $\ell = m$  et  $\phi(x_i) = x_i$ ,  $i = 1, \dots, m$ ), on se ramène à la contrainte de normalisation sur les poids  $\mathbf{w}^T \mathbf{w} = 1$  puisque la matrice  $\mathbf{B}$  sera une matrice identité.

On cherche à maximiser  $Var(t)$  sous la contrainte précédente :

$$\mathcal{L} = \mathbf{w}^T \mathbf{A} \mathbf{w} - \lambda \mathbf{w}^T \mathbf{B} \mathbf{w} \quad (4.41)$$

Ce qui revient à résoudre l'équation de vecteurs propres généralisés suivante :

$$\mathbf{A} \mathbf{w} = \lambda \mathbf{B} \mathbf{w} \quad (4.42)$$

où  $b_{ij} = \frac{\partial \phi_i}{\partial \mathbf{x}} \frac{\partial \phi_j}{\partial \mathbf{x}}$  et  $\frac{\partial \phi_i}{\partial \mathbf{x}} = \left( \frac{\partial \phi_i}{\partial x_1}, \dots, \frac{\partial \phi_i}{\partial x_m} \right)^T$

Nous pouvons écrire aussi :

$$\frac{1}{N} \sum_{k=1}^N \left( \frac{\partial t^T}{\partial \mathbf{x}} \frac{\partial t}{\partial \mathbf{x}} \right)_{\mathbf{x}=\mathbf{x}(k)} \equiv \mathbf{w}^T \mathbf{B} \mathbf{w} = 1 \quad (4.43)$$

Le modèle que nous venons de présenter définit une projection d'un espace de grande dimension vers un espace à une seule dimension. Étendons cette définition pour un espace de projection de dimension  $\ell > 1$  (fig 4.21). Ainsi, on cherche à calculer les poids de ce nouveau réseau à  $\ell$  composantes. La plus grande différence est par rapport à la contrainte de normalisation, où l'équation (4.40) est remplacée par :

$$\mathbf{W} \mathbf{B} \mathbf{W} = I_\ell \quad (4.44)$$

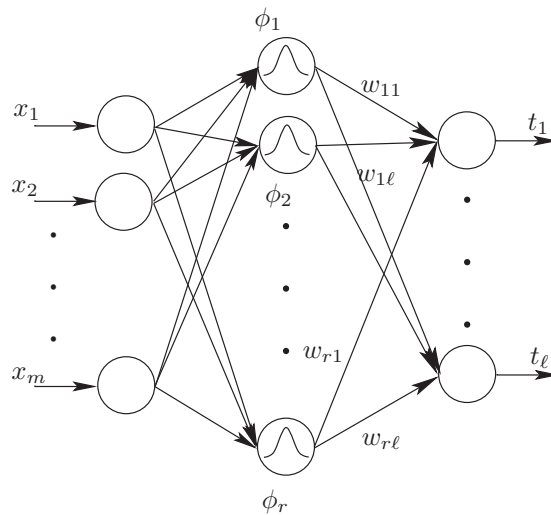


Figure 4.21 – Réseau RBF pour le calcul de  $\ell$  composantes principales

Ici,  $\mathbf{W} = [\mathbf{w}_1, \dots, \mathbf{w}_\ell]$  représente la matrice des vecteurs de poids. Ainsi la contrainte de normalisation peut être exprimée par :

$$\frac{1}{N} \sum_{k=1}^N \left( \frac{\partial t_i^T}{\partial \mathbf{x}} \frac{\partial t_j}{\partial \mathbf{x}} \right)_{\mathbf{x}=\mathbf{x}(k)} = [I_r]_{ij}, \quad 1 \leq i, j \leq \ell \quad (4.45)$$

La solution de  $\mathbf{W}$  peut être calculée en prenant les vecteurs propres correspondant aux  $\ell$  plus grandes valeurs propres à partir de la solution de la version matricielle de (4.42) :

$$\mathbf{A}\mathbf{W} = \lambda \mathbf{B}\mathbf{W} \quad (4.46)$$

Pour le calcul de la transformation inverse qui permettra d'estimer  $\hat{\mathbf{x}}$  à partir de la composante principale  $t$  donnée par le premier réseau RBF à trois couches, le deuxième réseau RBF (fig 4.20) est utilisé pour définir cette transformation inverse et qui est donnée par l'équation (4.36). Comme précédemment et pour simplifier en présentera le cas d'une seule composantes principale

Pour ce deuxième réseau, en supposant pour le moment que les centres et les dispersions des fonctions radiales sont optimisées, le problème de détermination des poids de la couche de sortie se ramène à un problème de regression linéaire.

Ainsi, pour la détermination des poids on doit minimiser l'erreur quadratique moyenne :

$$E = \frac{1}{N} \sum_{k=1}^N \|\mathbf{x}(k) - \mathcal{F}(t(k))\|^2 \quad (4.47)$$

où  $t(k) = \mathcal{G}(\mathbf{x}(k))$  et  $\mathcal{F}(t(k)) = \mathcal{F}(\mathcal{G}(\mathbf{x}(k)))$  est l'estimation de  $\mathbf{x}(k)$ .

En posant  $V_0 = [v_0 \ V]$ ,  $\psi_0(\cdot) \equiv 1$  et la matrice  $\Psi_0 = [1 \ \Psi]$  de dimension  $(N \times (r+1))$ , où la  $k^{eme}$  colonne de  $\Psi^T$  est  $\psi(t(k))$ ;  $\psi = (\psi_1, \dots, \psi_r)^T$  et  $X^T = (\mathbf{x}(1), \dots, \mathbf{x}(N))$ . On peut écrire :

$$E = \frac{1}{N} \sum_{k=1}^N \left\| \mathbf{x}(k) - \mathbf{V}_0^T \Psi_0(t) \right\|^2 \quad (4.48)$$

La solution de  $V_0$  qui minimise (4.48) est donnée par :

$$V_0^{opt} = \left( \Psi_0^T \Psi_0 \right)^{-1} \Psi_0^T X \quad (4.49)$$

à condition que l'inverse de  $(\Psi_0^T \Psi_0)$  existe.

Nous venons de présenter les deux approches permettant de calculer les poids des deux réseau. Cependant pour la détermination des centres et des dispersions des fonctions,  $\phi_i$  ( $i = 1, \dots, r$ ) et  $\psi_i$  ( $i = 1, \dots, r$ ), utilisées dans les neurones des couches cachées des deux réseaux ont doit minimiser l'erreur quadratique d'estimation donnée par :

$$E = \sum_{k=1}^N (\mathbf{x}_k - \hat{\mathbf{x}}_k)^2 \quad (4.50)$$

A partir de cette équation, on peut constater que l'apprentissage de ces deux réseau reste une tâche compliqué car on a besoin de  $\hat{\mathbf{x}}$  qui est donnée par le deuxième réseau. Ainsi, les deux réseaux ne sont pas totalement indépendants la procédure d'apprentissage doit mettre à jour les centres, les dispersions et les poids des deux réseaux.

Pour plus de simplification, nous proposons de combiner les réseau RBF à trois couches et les courbes principales.

#### 4.4.4.2 Approche combinant les courbes principales et les réseaux RBF

Cette approche cherche à simplifier le problème d'extraction des composantes principales non linéaires en utilisant deux réseaux RBF à trois couches en cascade dont les poids sont optimisés par un apprentissage supervisé. Ainsi, le problème est transformé en un problème de régression.

La sortie du réseau est alors une combinaison linéaire des variables transformées  $\phi(\mathbf{x})$ . Il faut noter que l'utilisation de l'algorithme des courbes principales de Verbeek [101] est plus adaptée à l'initialisation d'un réseau RBF car on a des centres et des dispersions définissant les fonctions noyau.

Considérons le réseau RBF illustré dans la figure (fig 4.19). L'objectif de ce réseau est de définir la transformation  $\mathcal{G} : \mathbb{R}^m \rightarrow \mathbb{R}$  permettant d'estimer la composante principale  $t$  qui est donnée par les équations (4.34) et (4.35).

où  $\{\phi_i, i = 1, \dots, r\}$  est une fonction non-linéaire de  $\mathbf{x} \in \mathbb{R}^m$  et  $\{w_i, i = 1, \dots, r\}$  représente l'ensemble des poids de la couche de sortie à déterminer. Les fonctions radiale de base  $\phi_i$  sont définis comme :

$$\phi_i(\mathbf{x}) = \exp\left(-\frac{\|\mathbf{x} - \mathbf{c}_i\|^2}{2\sigma_i^2}\right) \quad (4.51)$$

où  $\mathbf{c}_i$  et  $\sigma_i^2$  représentent, respectivement, les centres et les dispersions de ces fonctions.

Pour la détermination des centres, des dispersions et des poids du premier réseau RBF, l'erreur quadratique doit être minimisée :

$$E = \frac{1}{2} \sum_{k=1}^N (\mathbf{t}_k - \mathcal{G}(\mathbf{x}_k))^2 \quad (4.52)$$

Pour la détermination des règles d'adaptation des deux vecteurs de paramètres (centres et dispersions), on doit calculer le gradient de (4.52) par rapport aux centres  $\mathbf{c}_i$  et aux dispersions  $\sigma_i$ . Ainsi nous obtenons :

$$\frac{\partial E}{\partial \mathbf{c}_i} = \sum_{k=1}^N (\mathcal{G}(\mathbf{x}_k) - \mathbf{t}_k) \frac{\partial \mathcal{G}(\mathbf{x}_k)}{\partial \mathbf{c}_i} \quad (4.53)$$

avec  $\frac{\partial \mathcal{G}(\mathbf{x}_k)}{\partial \mathbf{c}_i} = w_i \frac{\mathbf{x}_k - \mathbf{c}_i}{\sigma_i^2} \phi_i(\mathbf{x}_k)$  et

$$\frac{\partial E}{\partial \sigma_i} = \sum_{k=1}^N (\mathcal{G}(\mathbf{x}_k) - \mathbf{t}_k) \frac{\partial \mathcal{G}(\mathbf{x}_k)}{\partial \sigma_i} \quad (4.54)$$

avec  $\frac{\partial \mathcal{G}(\mathbf{x}_k)}{\partial \sigma_i} = w_i \frac{\|\mathbf{c}_i - \mathbf{x}_k\|^2}{\sigma_i^3} \phi_i(\mathbf{x}_k)$ .

Ces deux dérivées sont utilisées pour la minimisation de (4.52) par la méthode du gradient et les poids  $w_i$  sont solution des moindres carrés minimisant (4.52) :

$$\mathbf{w} = (\Phi^T \Phi)^{-1} \Phi^T T \quad (4.55)$$

où la  $i^{eme}$  colonne de  $\Phi^T$  est  $\phi(\mathbf{x}_i) = (\phi_1(\mathbf{x}_i), \dots, \phi_r(\mathbf{x}_i))$ ,  $\mathbf{w}^T = (w_1, \dots, w_r)$  et la  $i^{eme}$  composante de  $T$  est  $\mathbf{t}_i$ .

L'apprentissage du deuxième réseau est identique au premier en ayant la composante  $t$  comme entrée et les variables originelles  $\mathbf{x}$  comme sorties. La solution est donnée par la même expression (4.49).

Cependant, l'inversion des matrices  $(\Phi^T \Phi)$  et  $(\Psi_0^T \Psi_0)$  peut être sujette à des problèmes numériques liés à son mauvais conditionnement. Pour éviter ce type de problème, des techniques de régularisation [29, 74] sont appliquées pour stabiliser la solution des moindres carrés.

Le résultat de l'application de l'approche proposée à notre exemple est illustré sur les figures (fig 4.22) et (fig 4.23).

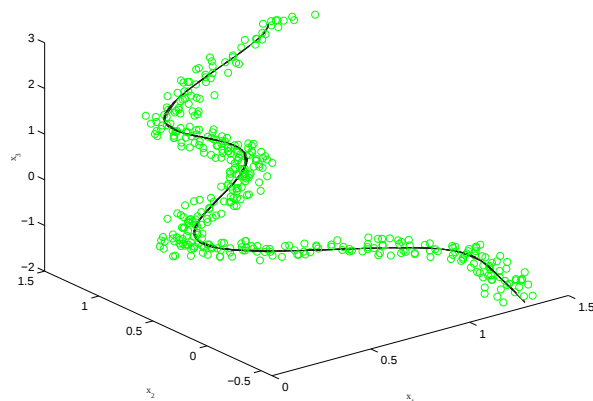


Figure 4.22 – Mesures et estimation de la courbe en combinant les courbes principales et deux réseaux RBF à trois couches (jeu d'identification)

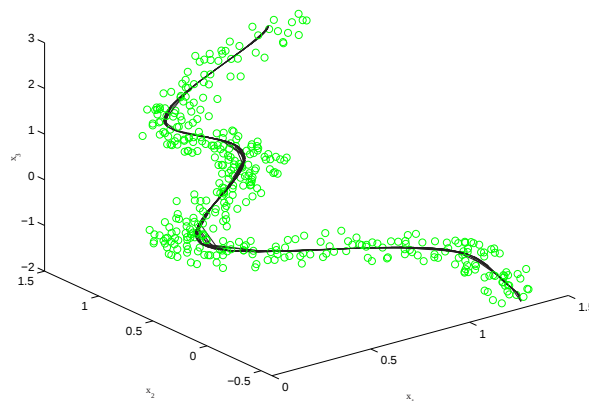


Figure 4.23 – Mesures et estimation de la courbe en combinant les courbes principales et deux réseaux RBF à trois couches (jeu de validation)

Le tableau (tab 4.1) présente les corrélations expliquées dans le cas des différentes approches.

|                             | ACP Linéaire | ACPNL 1 | ACPNL 2 | ACPNL 3 |
|-----------------------------|--------------|---------|---------|---------|
| Corrélations expliquée en % | 84 %         | 97 %    | 97 %    | 97 %    |

Table 4.1 – Corrélation expliquée par la première composante avec les quatre réalisations de l'ACP : ACP linéaire, ACNL 1 (réseau à cinq couches), ACPNL 2 (réseaux MLP à trois couches et courbes principales), ACPNL 3 (réseaux RBF et courbes principales)

Pour cet exemple, on a constaté que le fait d'utiliser les réseaux RBF pour construire le modèle ACPNL permet de simplifier le problème et de diminuer considérablement le temps de calcul tout en ayant la même qualité d'estimation des variables.

Dans notre travail nous n'avons pas traité le problème de l'optimisation de la structure des couches cachées. Cependant, nous nous sommes intéressés au problème de détermination du nombre de composantes à retenir dans le modèle ACPNL (nombre de neurones de la couche de sortie du premier réseau et qui représente le nombre de neurones de la couche d'entrée du deuxième réseau).

#### 4.4.4.3 Détermination du nombre de composantes non linéaires

Webb [105] propose d'utiliser un indice classique pour la détermination du nombre de composantes non linéaires :

$$\varepsilon = \left( \frac{\|\hat{X} - X\|^2}{\|X - \bar{X}\|^2} \right)^{\frac{1}{2}} \quad (4.56)$$

où  $\hat{X}$  est l'estimation de  $X$  par le modèle neuronal et  $\bar{X}$  est une matrice dont les lignes représentent le vecteur de moyenne de la matrice  $X$ .

Dans le cas d'une ACP linéaire utilisant les  $\ell$  premières composantes :

$$\varepsilon = \left( \frac{\sum_{i=\ell+1}^m \lambda_i}{\sum_{i=1}^m \lambda_i} \right)^{\frac{1}{2}} \quad (4.57)$$

ce qui donne la relation entre  $\eta$  et l'erreur normalisée  $\varepsilon$  :

$$\varepsilon^2 = 1 - \eta \quad (4.58)$$

où  $\eta = \frac{\sum_{i=1}^{\ell} \lambda_i}{\sum_{i=1}^m \lambda_i}$  donne la fraction de la variance retenue dans un espace de dimension  $\ell$  défini par les vecteurs propres associés à ces valeurs propres.

Dans le cas non linéaire, on va augmenter le nombre  $\ell$  progressivement et on surveille la valeur de  $\varepsilon$ . Cependant, nous avons déjà vu, dans le cas linéaire, que l'utilisation d'un tel indice tend à considérer un nombre excessif de composantes. Pour cette raison, nous proposons d'étendre le principe de la variance non reconstruite dans le cas non linéaire.

A partir des données représentées par le vecteur  $\mathbf{x}$ , on peut calculer le vecteur des composantes principales  $\mathbf{t}$  par :

$$\mathbf{t} = \sum_{j=1}^r w_j \phi(y_j) \quad (4.59)$$

où  $y_j = -\frac{1}{2} \left\| \frac{\mathbf{x} - \mathbf{c}_j}{\sigma_j} \right\|^2$ . De plus, on peut obtenir l'estimation de  $\mathbf{x}$  par :

$$\hat{\mathbf{x}} = V^T \Psi(\mathbf{t}) = \begin{bmatrix} v_{01} & v_{02} & . & . & . & v_{0r} \\ v_{11} & v_{12} & . & . & . & v_{1r} \\ v_{21} & . & . & . & . & v_{2r} \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ . & . & . & . & . & . \\ v_{m1} & v_{m2} & . & . & . & v_{mr} \end{bmatrix} \begin{bmatrix} \psi_0(\mathbf{t}) \\ \psi_1(\mathbf{t}) \\ \psi_2(\mathbf{t}) \\ . \\ . \\ . \\ \psi_r(\mathbf{t}) \end{bmatrix} \quad (4.60)$$

A partir de ces deux expressions, on peut calculer l'erreur d'estimation  $\mathbf{e} = \mathbf{x} - \hat{\mathbf{x}}$ , puis le *SPE* :

$$SPE = \mathbf{e}^T \mathbf{e} = \|\mathbf{x} - \hat{\mathbf{x}}\|^2 \quad (4.61)$$

La reconstruction de la  $i^{eme}$  variable notée  $z_i$  peut être effectuée, comme dans le cas linéaire, par minimisation du *SPE* par rapport à cette variable.

$$\frac{\partial SPE}{\partial z_i} = \left( \xi_i^T - \frac{\partial \psi^T(\mathbf{t})}{\partial z_i} V \right) (\mathbf{x} - \hat{\mathbf{x}}) = 0 \quad (4.62)$$

On a les résultats intermédiaires suivants :

$$\frac{\partial \psi^T(\mathbf{t})}{\partial z_i} = \frac{\partial \mathbf{t}}{\partial z_i} \frac{\partial \psi^T(\mathbf{t})}{\partial \mathbf{t}} \quad (4.63)$$

$$\frac{\partial \mathbf{t}}{\partial z_i} = \sum_{j=1}^r w_j \frac{\partial \phi(y_j)}{\partial z_i} \quad (4.64)$$

$$\frac{\partial \phi(y_j)}{\partial z_i} = \frac{\partial y_j}{\partial z_i} \frac{\partial \phi(y_j)}{\partial y_j} \quad (4.65)$$

$$\frac{\partial y_j}{\partial z_i} = -\frac{1}{\sigma_j} \xi_i^T \left( \frac{\mathbf{x} - c_j}{\sigma_j} \right) \quad (4.66)$$

$$\frac{\partial SPE}{\partial z_i} = \left[ \xi_i^T + \sum_{j=1}^r \left\{ w_j \frac{1}{\sigma_j} \xi_i^T \left( \frac{\mathbf{x} - c_j}{\sigma_j} \right) \frac{\partial \phi(y_j)}{\partial y_j} \right\} \frac{\partial \psi^T(\mathbf{t})}{\partial \mathbf{t}} V \right] (\mathbf{x} - \hat{\mathbf{x}}) = 0$$

$$\xi_i^T \mathbf{x} - \xi_i^T \hat{\mathbf{x}} + \sum_{j=1}^r \left\{ w_j \frac{1}{\sigma_j} \xi_i^T \left( \frac{\mathbf{x} - c_j}{\sigma_j} \right) \frac{\partial \phi(y_j)}{\partial y_j} \right\} \frac{\partial \psi^T(\mathbf{t})}{\partial \mathbf{t}} V (\mathbf{x} - \hat{\mathbf{x}}) = 0$$

Or  $\hat{\mathbf{x}} = V^T \psi(\mathbf{t})$  et  $\xi_i^T \mathbf{x} = z_i$ .

$$z_i - \xi_i^T V^T \psi(\mathbf{t}) + \sum_{j=1}^r \left[ w_j \left( \frac{z_i - c_{ij}}{\sigma_j^2} \right) \frac{\partial \phi_j}{\partial y_j} \right] \frac{\partial \psi^T}{\partial \mathbf{t}} V (\mathbf{x} - V^T \psi(\mathbf{t})) = 0$$

$$z_i = \xi_i^T V^T \psi(\mathbf{t}) - \sum_{j=1}^r \left[ w_j \left( \frac{z_i - c_{ij}}{\sigma_j^2} \right) \frac{\partial \phi_j}{\partial y_j} \right] \frac{\partial \psi^T}{\partial \mathbf{t}} V (\mathbf{x} - V^T \psi(\mathbf{t})) \quad (4.67)$$

Dans le cas linéaire nous avons :  $\psi(\mathbf{t}) = \mathbf{t}$ ,  $\phi(y) = y$ ,  $c_j = 0$ ,  $\sigma_j = 1$ ,  $V^T = \hat{P}$ ,  $W^T = \hat{P}^T$  et ainsi nous avons :  $\mathbf{t} = \hat{P}^T \mathbf{x}$  et  $\hat{\mathbf{x}} = \hat{P} \mathbf{t} = \hat{P} \hat{P}^T \mathbf{x} = \hat{C} \mathbf{x}$ . En remplaçant les nouvelles variables dans l'expression de  $z_i$  dans le cas non linéaire, nous obtenons :

$$z_i = \xi_i^T \hat{P} \hat{P}^T \mathbf{x} - \xi_i^T \hat{P}^T \mathbf{x} \hat{P}^T (\hat{P} \hat{P}^T \mathbf{x} - \mathbf{x}) = \xi_i^T \hat{P} \hat{P}^T \mathbf{x} \quad (4.68)$$

Cette expression est la même que celle de l'équation (2.41).

La valeur de la variable à reconstruire peut se calculer grâce à la méthode des valeurs manquantes comme dans le cas linéaire par la recherche de la valeur de  $\mathbf{t}$ , notée  $\hat{\mathbf{t}}$ , qui minimisera la quantité suivante :

$$\hat{\mathbf{t}} = \arg \min_{\mathbf{t}} \left\| \mathbf{x}^{(i)} - V_i^T \psi(\mathbf{t}) \right\|^2 \quad (4.69)$$

où  $\mathbf{x}^{(i)}$  représente le vecteur des mesures sans la  $i^{eme}$  variables et  $V_i$  est la matrice des poids dont la  $i^{eme}$  colonne est éliminée.

$$\frac{\partial SPE^{(i)}}{\partial \mathbf{t}} = \frac{\partial \left( \left\| \mathbf{x}^{(i)} - V_i^T \psi(\mathbf{t}) \right\|^2 \right)}{\partial \mathbf{t}} \quad (4.70)$$

$$\frac{\partial SPE^{(i)}}{\partial \mathbf{t}} = 2 \frac{\partial \psi^T}{\partial \mathbf{t}} V_i \left( \mathbf{x}^{(i)} - V_i^T \psi(\mathbf{t}) \right) = 0 \quad (4.71)$$

A partir de l'équation (4.71), l'expression de  $z_i$  de l'équation (4.67) peut s'écrire sous la forme :

$$z_i = \xi_i^T V^T \psi(\mathbf{t}) - \sum_{j=1}^r \left[ w_j \left( \frac{z_i - c_{ij}}{\sigma_j^2} \right) \frac{\partial \phi_j}{\partial y_j} \right] \frac{\partial \psi^T}{\partial \mathbf{t}} v_i \left( z_i - v_i^T \psi(\mathbf{t}) \right) \quad (4.72)$$

car :

$$\frac{\partial \psi^T}{\partial \mathbf{t}} V \left( \mathbf{x} - V^T \psi(\mathbf{t}) \right) = \frac{\partial \psi^T}{\partial \mathbf{t}} V_i \left( x^{(i)} - V_i^T \psi(\mathbf{t}) \right) + \frac{\partial \psi^T}{\partial \mathbf{t}} v_i \left( z_i - v_i^T \psi(\mathbf{t}) \right) \quad (4.73)$$

où  $v_i^T$  est la  $i^{ieme}$  colonne de la matrice  $V$ .

Puisque la reconstruction  $z_i$  est obtenue par projection répétée via le modèle non linéaire, à la convergence  $z_i = v_i^T \psi(\mathbf{t})$  et ainsi le second terme de l'équation (4.72) s'annulera. Si on note  $\hat{\mathbf{t}}$  la valeur de  $\mathbf{t}$  obtenue à la convergence par  $\hat{\mathbf{t}}$ , alors l'expression de  $z_i$  peut s'écrire sous la forme :

$$z_i = \xi_i^T V^T \psi(\hat{\mathbf{t}}) \quad (4.74)$$

Ainsi nous proposons une méthode de reconstruction dans le cas non linéaire pour la détermination du nombre de composantes à retenir dans le modèle ACPNL et qui peut être utiliser pour la localisation de défauts. Cette méthode utilise l'équation (4.72) pour la reconstruction des variables :

$$z_i^{(iter)} = \xi_i^T V^T \psi(\mathbf{t}) - \sum_{j=1}^r \left[ w_j \left( \frac{z_i^{(iter-1)} - c_{ij}}{\sigma_j^2} \right) \frac{\partial \phi_j}{\partial y_j} \right] \frac{\partial \psi^T}{\partial \mathbf{t}} v_i \left( z_i^{(iter-1)} - v_i^T \psi(\mathbf{t}) \right) \quad (4.75)$$

Comme c'est une méthode itérative, il faut choisir une initialisation pour  $z_i$ . Nous avons choisi de poser comme valeur initiale ( $iter = 0$ )  $z_i^{(0)} = x_i$ . Après convergence  $z_i$  est la valeur reconstruite de  $x_i$ . Ainsi, la variance non reconstruite de la  $i^{eme}$  variable, en utilisant un modèle avec  $\ell$  composantes, est définie comme dans le cas linéaire par :

$$\rho_i(\ell) = var\{x_i - z_i\} = \mathcal{E} \left\{ (x_i - z_i)^2 \right\} \quad (4.76)$$

et le critère à minimiser par rapport à  $\ell$  est donné par :

$$\min_{\ell} \sum_{i=1}^m \rho_i(\ell) \quad (4.77)$$

Pour illustrer cette approche, nous proposons l'exemple à deux variables suivant :

$$\begin{aligned} x_1 &= a^2 + \epsilon_1 \\ x_2 &= a + \epsilon_2 \end{aligned} \quad (4.78)$$

où  $a \in [-1, 1]$  est une variable aléatoire uniforme et  $\epsilon_i$  un bruit aléatoire uniformément distribué dans l'intervalle  $[-0.1, 0.1]$ .

Le modèle ACPNL à deux réseaux RBF à trois couches est utilisé. Chaque réseau a trois neurones dans sa couche cachée. La variance non reconstruite calculée (4.77), pour les différents nombre de composantes  $\ell = 1, 2, 3$ , est représentée dans le tableau (tab 5.3).

| $var(1)$ | $var(2)$ | $var(3)$ |
|----------|----------|----------|
| 0.0190   | 0.0194   | 0.0198   |

Table 4.2 – Variance non reconstruite en fonction du nombre de composantes

Ainsi, le nombre de composantes à retenir dans le modèle est  $\ell = 1$ . L'estimation fournie par le modèle ACPNL à une composante est représenté sur la figure (fig 4.24).

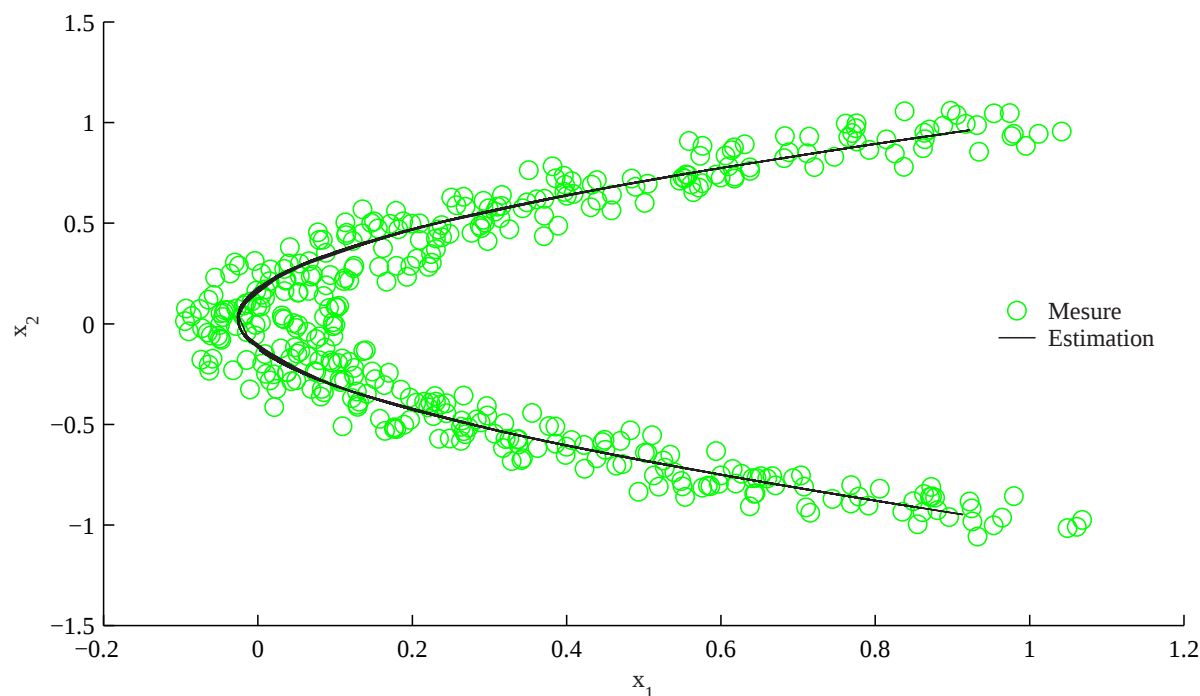


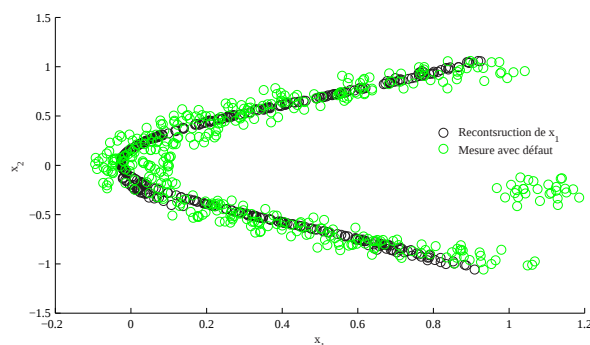
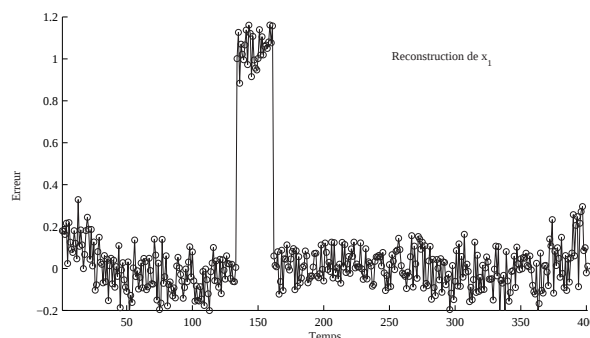
Figure 4.24 – Mesure et estimation en utilisant une seule composante

Dans la suite nous allons simuler un biais d'une amplitude égale à 1 sur la première variable entre les instants 134 et 161. Les figures (fig 4.25) et (fig 4.26) représentent respectivement les mesures avec le défaut et la reconstruction de la première variable ainsi que l'erreur de reconstruction.

A partir de ces dernières figures, la reconstruction de la variable  $x_1$  est très satisfaisante et l'amplitude du défaut peut être estimée en calculant l'erreur de reconstruction (fig 4.25).

Il faut noter que l'algorithme converge assez rapidement, toutefois, il reste à trouver les conditions de convergence.

Une fois le modèle ACPNL déterminé, les résidus peuvent être générés pour effectuer le diagnostic de fonctionnement du processus.

Figure 4.25 – Mesures et reconstruction de la courbe de l'exemple 2 avec un défaut sur  $x_1$ Figure 4.26 – Erreur de reconstruction de la variable  $x_1$ 

## 4.5 Conclusion

Dans ce chapitre nous avons présenté le principe de l'analyse en composantes principales non linéaire. Ainsi, plusieurs approches utilisées pour le calcul du modèle ACPNL sont présentées. Dans un premier temps, nous avons présenté le principe de l'approche des courbes principales. Cette approche est une généralisation non paramétrique de l'analyse en composantes principales dans le cas non-linéaire. Cependant, elle ne permet pas d'avoir un modèle de représentation et donc elle n'est pas intéressante pour des objectifs de diagnostic.

Les approches neuronales semblent plus intéressantes pour la modélisation des systèmes non linéaires. Ainsi, plusieurs approches sont présentées utilisant différentes structures des réseaux de neurones. Le réseau le plus utilisé pour le calcul de l'ACPNL est un réseau à cinq couches. Vu la taille du réseau, le nombre de paramètres à optimiser est considérable et donc l'apprentissage est plus difficile. Pour simplifier l'apprentissage, d'autres alternatives ont été proposées. Ainsi, on trouve le réseau à trois couches (IT network) dont la procédure d'optimisation cherche à optimiser les entrées de ce réseau (représentant les composantes principales) en plus de l'optimisation des poids. L'autre solution de ce problème consiste à combiner les courbes principales et les réseaux de neurones à trois couches. Ainsi, le modèle ACPNL est obtenu par apprentissage supervisé de deux réseaux de neurones à trois couches. Cependant, comme l'apprentissage de ces réseaux pose souvent énormément de problèmes de convergence et d'initialisation, les réseaux RBF semblent être plus intéressants à utiliser.

Ainsi, nous proposons une méthode pour le calcul du modèle ACPNL en combinant des réseaux RBF à trois couches et l'approche des courbes principales. Ce qui permet d'obtenir un modèle ACPNL dont la solution des poids est donnée par une estimation des moindres carrés pour les deux réseaux. Ainsi, le temps de calcul est considérablement réduit par rapport aux approches utilisant les réseaux de neurones. De plus, pour la détermination du nombre de composantes principales, les réseaux RBF offrent la possibilité d'étendre le critère de la variance non reconstruite pour la sélection du nombre de composantes dans le cas non-linéaire. Ainsi, nous avons proposé une méthode pour la détermination de ce nombre en exploitant le critère de minimisation de la variance non reconstruite développé dans le cas linéaire. Les résultats obtenus sont très satisfaisants. Il nous reste à trouver les conditions de convergence de la méthode proposée. Les résultats obtenus sont encourageants pour appliquer les procédures proposées sur des systèmes plus compliqués et en optimisant la structure des réseaux RBF utilisés.



# 5

## Détection et localisation de défauts capteurs d'un réseau de surveillance de la qualité de l'air

### Sommaire

---

|            |                                                                               |            |
|------------|-------------------------------------------------------------------------------|------------|
| <b>5.1</b> | <b>Introduction . . . . .</b>                                                 | <b>131</b> |
| <b>5.2</b> | <b>Description du phénomène . . . . .</b>                                     | <b>132</b> |
| <b>5.3</b> | <b>Description du réseau de surveillance de la qualité de l'air . . . . .</b> | <b>133</b> |
| <b>5.4</b> | <b>Prétraitement des données . . . . .</b>                                    | <b>136</b> |
| <b>5.5</b> | <b>Application de l'ACP linéaire . . . . .</b>                                | <b>139</b> |
| <b>5.6</b> | <b>Conclusion . . . . .</b>                                                   | <b>151</b> |

---

### 5.1 Introduction

Beaucoup d'activités humaines produisent des polluants primaires comme les oxydes d'azote ( $\text{NO}_2$  et  $\text{NO}$ ), le dioxyde de soufre et les composés organiques volatiles (COV) qui forment dans la basse atmosphère, par des réactions chimiques ou photochimiques, des polluants secondaires comme l'ozone.

Un certain nombre de ces polluants sont susceptibles de poser des problèmes pour la santé humaine et les systèmes écologiques. C'est pourquoi une directive européenne a défini des normes de qualité de l'air afin de protéger la santé humaine. Ainsi, les valeurs seuils suivantes ont été fixées pour la concentration d'ozone :

- $360\mu\text{g}/\text{m}^3$  (valeur moyenne sur une heure) : seuil d'alerte de la population,
- $180\mu\text{g}/\text{m}^3$  (valeur moyenne sur une heure) : seuil d'information de la population,
- $110\mu\text{g}/\text{m}^3$  (valeur moyenne sur une heure) : seuil de protection pour la santé.

et des seuils ont été également étendus aux  $\text{NO}_x$ .

La surveillance de la qualité de l'air est effectuée par les réseaux de mesures. Leurs missions sont : la production de données (mesures de concentration de polluants et d'un ensemble de paramètres météorologiques liés aux événements de pollution) comprenant la gestion des réseaux

de mesures, la diffusion des données pour l'information permanente de la population et des services publics en cas de dépassement des seuils fixés.

Des données fausses entraîneront des décisions erronées (mesures temporaires de réduction des émissions). Ainsi, la validité des données et la crédibilité des informations fournies sont donc essentielles du point de vue économique, sanitaire et écologique, social, scientifique et technique. Pour garantir la pertinence des informations fournies ou des décisions prises, la mise en oeuvre de procédures de diagnostic de fonctionnement de capteurs et de validation de données est indispensable avant toute utilisation des mesures par un opérateur ou un système de traitement de l'information.

La validation de données a pour but de détecter des anomalies de fonctionnement des capteurs principalement ceux de la concentration d'ozone et des polluants primaires comme les oxydes d'azote NO et NO<sub>2</sub>.

Ainsi, on est typiquement en face d'un problème de diagnostic qui nécessite les étapes suivantes :

- détection de capteurs défaillants,
- localisation du ou des capteurs défaillants,
- identification du type de défauts (amplitude, durée),
- proposition d'une valeur de remplacement (utile pour les traitements ultérieurs).

Le réseau comporte de nombreux sites de mesures réparties sur une zone géographique importante. De plus on a de nombreux appareils de mesures sur chacun des sites.

Actuellement, la validation est effectuée manuellement par un opérateur et donc trop subjective. De plus, c'est une tâche fastidieuse vu le nombre de capteurs à surveiller.

De plus, comme il y a beaucoup de capteurs et qu'apparemment il y a des corrélations importantes entre ces derniers, la recherche de modèles sous forme entrée/sortie pour chaque variable devient dans ce cas une tâche très difficile d'où la nécessité d'utiliser une méthode de modélisation plus adéquate permettant de modéliser le comportement de l'ensemble des variables. Une des méthodes les plus utilisées dans ce cadre est l'analyse en composantes principales.

Ce chapitre est consacré à l'application de l'analyse en composantes principales pour la détection et la localisation de défauts de capteurs d'un réseau de surveillance de la qualité de l'air en Lorraine. Il faut rappeler ici que l'étude présentée est une étude de faisabilité et donc on s'est limité à des courtes périodes de mesures. AIRLOR n'a pas d'exemple de défauts à nous montrer. Mais pour leur démarche "qualité", ils ont besoin d'un système permettant de certifier la qualité de leurs mesures. On a donc décidé de simuler des défauts.

Après une description du phénomène de la pollution atmosphérique, les données traitées seront présentées. Ainsi, un modèle ACP sera calculé pour la modélisation du comportement du fonctionnement normal du réseau. Une fois le modèle ACP calculé on passe à la phase de génération de résidus pour la détection et la localisation de défauts des différents capteurs considérés dans cette étude. Enfin, on présentera l'utilisation des méthodes développées dans la partie théorique avec succès sur au moins une partie de la base de données.

## 5.2 Description du phénomène

L'ozone est un constituant de la troposphère. Il n'est pas émis directement dans l'atmosphère ; il est le produit de processus physiques et chimiques de transformation de composés primaires. Ces processus sont complexes par le nombre de réactions chimiques mises en jeu, par les vitesses très différentes qu'elles peuvent avoir suivant les concentrations des polluants de la zone géographique considérée.

Dans une atmosphère non polluée, l'ozone résulte principalement de la seule réaction de combinaison d'un atome d'oxygène O avec l'oxygène de l'air O<sub>2</sub> en présence d'un corps stabilisant M (fig 5.1). L'atome d'oxygène nécessaire à cette réaction est obtenu par photodissociation du dioxyde d'azote NO<sub>2</sub> en NO et O. Mais la molécule de NO ainsi formée est oxydée rapidement par l'ozone pour reformer le NO<sub>2</sub>. Il s'établit un cycle appelé cycle de Chapman [1]. Un régime stationnaire s'établit, caractérisé par une concentration d'ozone plus ou moins constante qui dépend des concentrations de NO et de NO<sub>2</sub> et des vitesses des trois réactions.

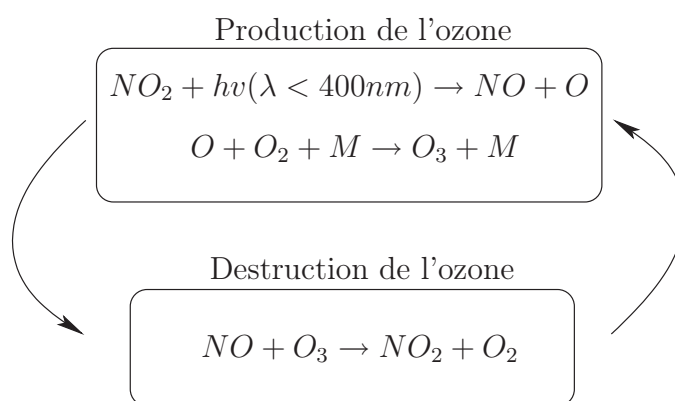


Figure 5.1 – Illustration du cycle de Chapman.

Dès lors, une augmentation de la concentration d'ozone est due à une transformation du NO en NO<sub>2</sub> sans consommation de molécules d'ozone. En atmosphère polluée, l'action de produits comme les composés organiques volatiles COV (hydrocarbures et composés oxygénés) et l'aérosol urbain perturbent le cycle de Chapman en offrant des voies d'oxydation des NO autres que celle de l'ozone. L'action de ces composés conduit, à travers une série de réactions chimiques complexes, à une oxydation de NO en NO<sub>2</sub> sans destruction de l'ozone. Les molécules de NO<sub>2</sub> formées sont ensuite dissociées sous l'action de la lumière. Les atomes d'oxygène O qui en résultent alimentent le processus de production de l'ozone. Par conséquent, on comprend aisément l'augmentation de la teneur en ozone dès que le rayonnement solaire est suffisamment intense. Une partie des NO participe à la destruction de l'ozone alors que l'autre partie se combine de nouveau beaucoup plus rapidement avec les COV pour produire du NO<sub>2</sub>. Il faut noter que l'ozone O<sub>3</sub> est moins localisé et donc plus facile à modéliser. De plus, comme les COV et les oxydes d'azote NO et NO<sub>2</sub> sont des polluants émis et liés à l'activité humaine (transport, chauffage collectif, industries,...) ils sont donc plus localisés que l'ozone et donc plus difficiles à relier les mesures des différents sites.

### 5.3 Description du réseau de surveillance de la qualité de l'air

Lors de l'étude, le réseau de surveillance AIRLOR se compose de onze stations placées dans des sites ruraux, périurbains et urbains (fig 5.2). Chaque station de surveillance se compose d'un ensemble de capteurs de mesure des concentrations des polluants suivants : le monoxyde de carbone CO, les oxydes d'azote (NO et NO<sub>2</sub>) qui sont mesurées par le même analyseur, le dioxyde de soufre SO<sub>2</sub> et l'ozone O<sub>3</sub>. Certaines stations mesurent en plus certains paramètres météorologiques.

**Remarque :** les concentrations des polluants sont en  $\mu g/m^3$ .

Les grandeurs météorologiques mesurées sont les suivantes :

- température ( $^{\circ}\text{C}$ ),
- humidité relative (%),
- rayonnement solaire global ( $\text{W}/\text{m}^2$ ),
- pression atmosphérique (hPa),
- vitesse du vent (m/s),
- direction du vent (degré).

Une station de mesure est un local dans lequel se trouvent des analyseurs. L'air extérieur est pompé et amené jusqu'à l'analyseur qui va mesurer sa teneur en un ou plusieurs polluants spécifiques. L'air est prélevé à 3 mètres environ. La mesure se fait en continu tous les quarts d'heure, 24 heures sur 24 et toute l'année. Les stations de mesures, réparties sur différents secteurs géographique, sont reliées à un ordinateur central par ligne téléphonique. Les données sont ainsi télétransmises tous les jours. Elles sont validées manuellement, traitées statistiquement et analysées avant d'être transmises aux médias et à une banque de données nationale (fig 5.3).

Figure 5.2 – Implantation des capteurs, réseau AIRLOR .

Nous disposons des mesures pour les sites de mesures du réseau AIRLOR. Les dates indiquées sur le tableau (tab 5.1) correspondent aux jours pour lesquels les premières mesures sont disponibles.

Les figures (fig 5.4), (fig 5.5) et (fig 5.6) présentent l'évolution des concentrations de  $\text{O}_3$ ,  $\text{NO}_2$  et  $\text{NO}$  pour les trois sites Brabois, Nancy Kennedy (Dan) et Fléville.

Sur ces figures, il peut être remarqué que la concentration d'ozone est modulée par un cycle jour-nuit et présente une évolution quotidienne sous la forme d'une courbe en cloche. Faible la nuit, le niveau d'ozone augmente progressivement en début de journée pour atteindre ses valeurs maximales dans l'après-midi. Il descend ensuite à une valeur nocturne faible et comparable à celle de la veille.

Ainsi, on remarque qu'il y a des valeurs manquantes pour le site de Brabois par exemple pour les trois signaux  $\text{O}_3$ ,  $\text{NO}$  et  $\text{NO}_2$  (panne de batterie par exemple). On constate également que le  $\text{O}_3$  de Brabois et celui de Fléville sont comparables en évolution et en ordre de grandeur alors que le  $\text{O}_3$  de Dan est comparable uniquement en évolution par rapport aux deux premiers (fig 5.4). Ainsi, il est possible de trouver des modèles reliant ces différentes grandeurs. Les  $\text{NO}_2$

Figure 5.3 – Schéma du système d'acquisition et d'archivage des données caractérisant la qualité de l'air.

| Station          | $NO_x$   |          | $O_3$    |          |          |
|------------------|----------|----------|----------|----------|----------|
| Bar le Duc       | 05/01/95 | 05/01/95 | 05/01/95 | 05/01/95 | 05/01/95 |
| Brabois          | 03/07/94 | 03/07/94 | -        | 03/07/94 | 03/07/94 |
| Nancy Centre DAN | 00/07/94 | 00/07/94 | 00/07/94 | 00/07/94 | -        |
| Fléville         | 27/01/95 | 27/01/95 | -        | 27/01/95 | -        |
| Tomblaine        | 12/12/95 | 12/12/95 | -        | 12/12/95 | -        |
| St Nicolas       | 04/07/95 | 04/07/95 | -        | 04/07/95 | 04/07/95 |
| Lunéville        | 11/12/95 | 11/12/95 | -        | 11/12/95 | -        |
| Epinal           | 24/03/95 | 24/03/95 | 24/03/95 | 24/03/95 | 24/03/95 |

Table 5.1 – Sites de mesures disponibles du réseau AIRLOR

par contre sont très différents pour les trois sites (fig 5.5) alors que pour le NO les mesures sur les trois sites sont un peu plus comparables (fig 5.6) que les NO<sub>2</sub>. D'où le problème de modélisation de ces deux derniers polluants.

En se référant au tracé comparatif de l'ozone et des  $\text{NO}_x$ , on peut constater que les concentrations de  $\text{NO}_x$  et d'ozone varient en raison inverse. La nuit, les  $\text{NO}_x$  sont présents alors que l'ozone est piégé. En début de journée, quand le niveau d'ozone monte, il y a décroissance de celui des  $\text{NO}_x$ . La relation inverse s'établit en fin de journée au moment de la baisse de la concentration d'ozone (fig 5.7) et (fig 5.8).

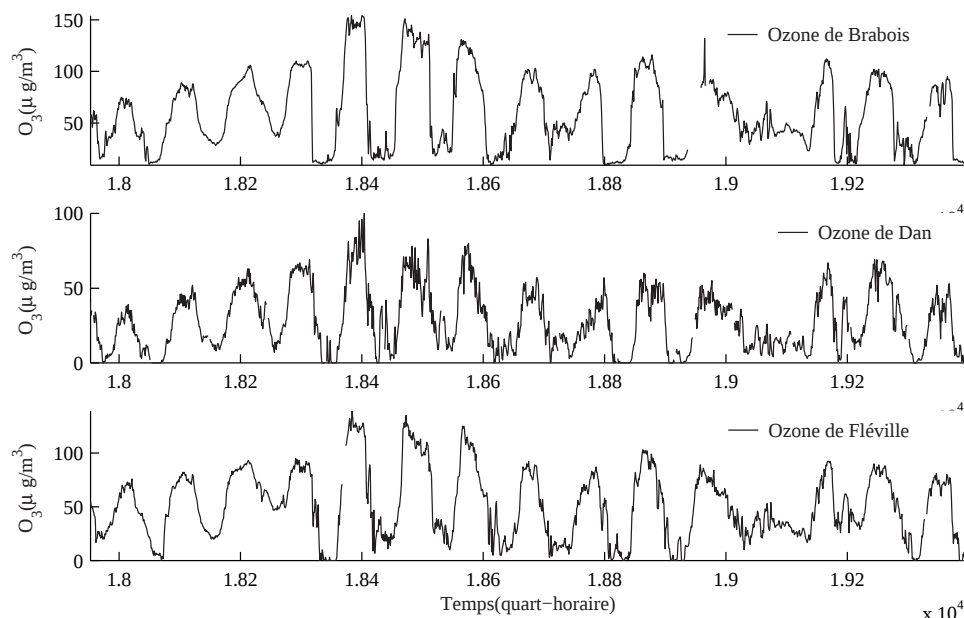


Figure 5.4 – Concentration d'ozone des sites de Brabois, Dan et Fléville sur une période de quelques jours

## 5.4 Prétraitement des données

Les archives de mesures mises à disposition par les réseaux contiennent des séries de valeurs manquantes. Celles-ci peuvent avoir diverses causes :

- panne de la chaîne d'instrumentation (batterie, capteur, ...),
- maintenance des capteurs,
- mesure refusée par le réseau (mesure aberrante).

Ces valeurs manquantes sont autant d'anomalies qui posent problèmes lors des traitements ultérieurs, il est donc important de les éliminer. Selon la durée de l'absence de mesure, il existe plusieurs solutions :

- remplacement des valeurs manquantes par simple interpolation linéaire, dans le cas de séries de quelques points de mesures,
- remplacement par estimation ou par reconstruction à partir des mesures des autres variables,
- dans le cas où les valeurs manquantes présentent une série de mesures sur plusieurs jours ou sur plusieurs variables (exemple du  $\text{NO}$  et  $\text{NO}_2$  sur les figures précédentes), il est préférable d'éliminer l'ensemble des données correspondant à cette période pour toutes les variables.

D'après l'analyse que nous avons effectuée [30] pour la modélisation des concentrations d'ozone, nous avons constaté que l'utilisation d'un modèle linéaire statique entre la mesure de

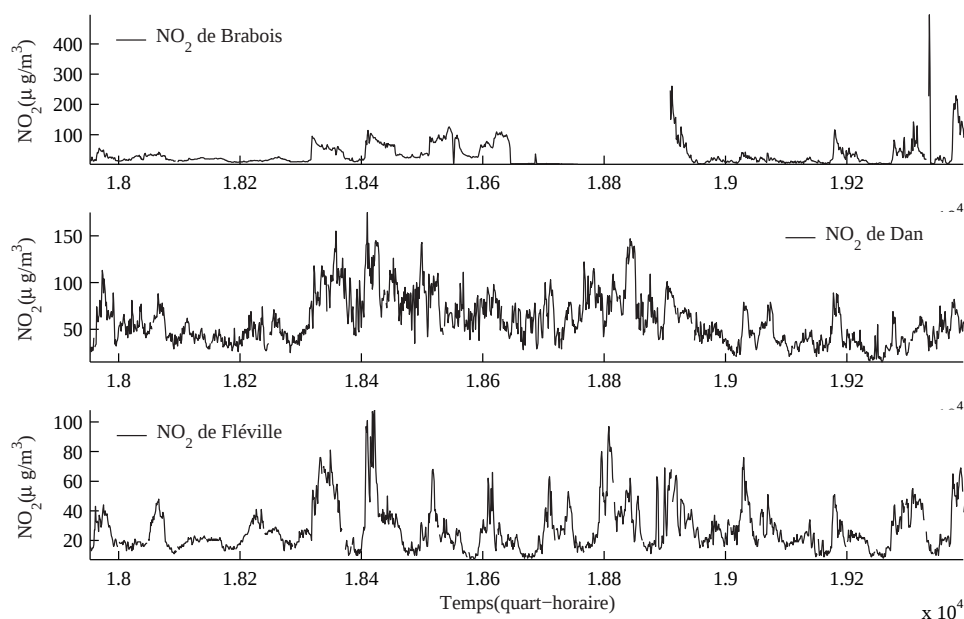


Figure 5.5 – Concentration de  $NO_2$  des sites de Brabois, Dan et Fléville sur une période de quelques jours

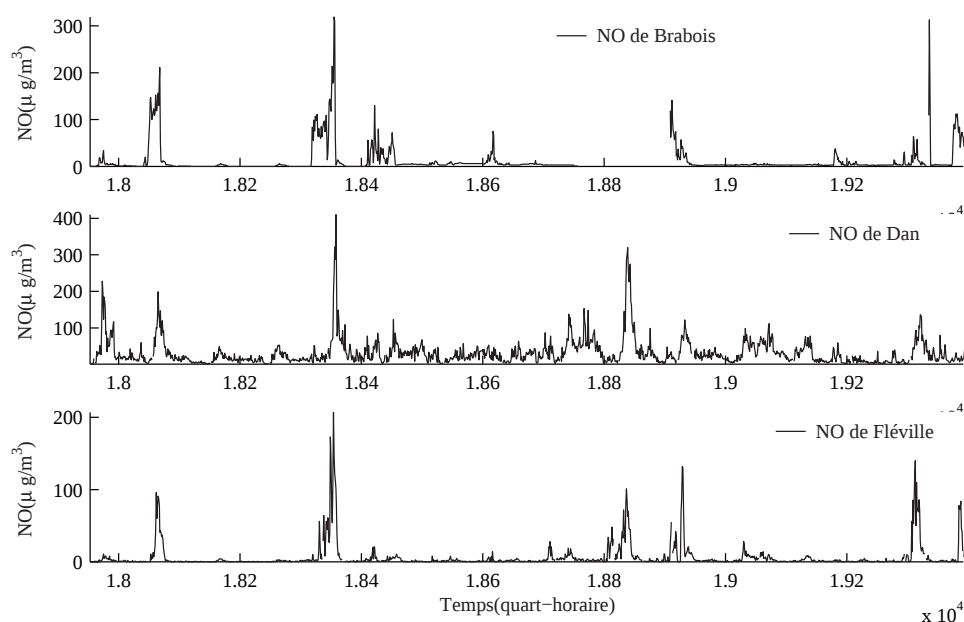
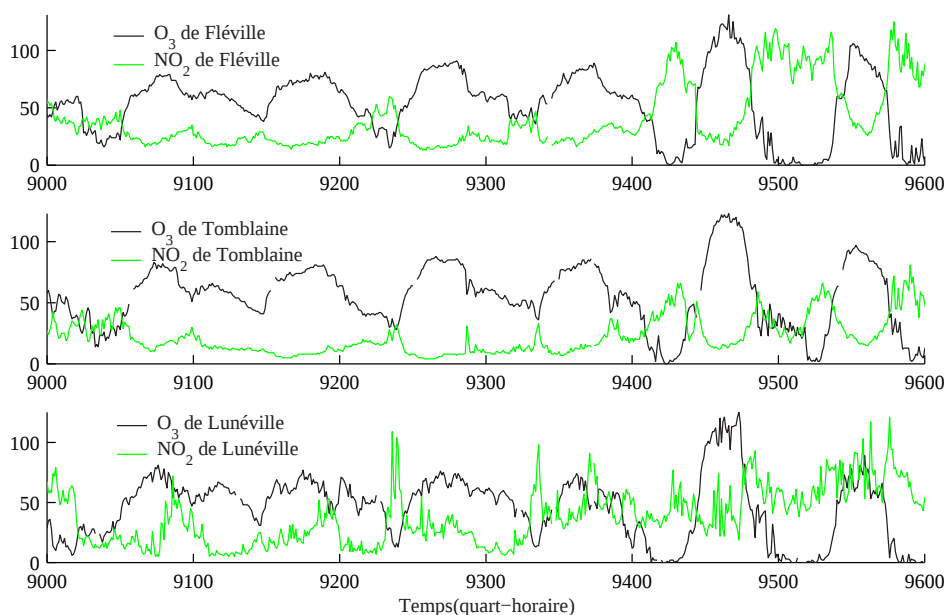
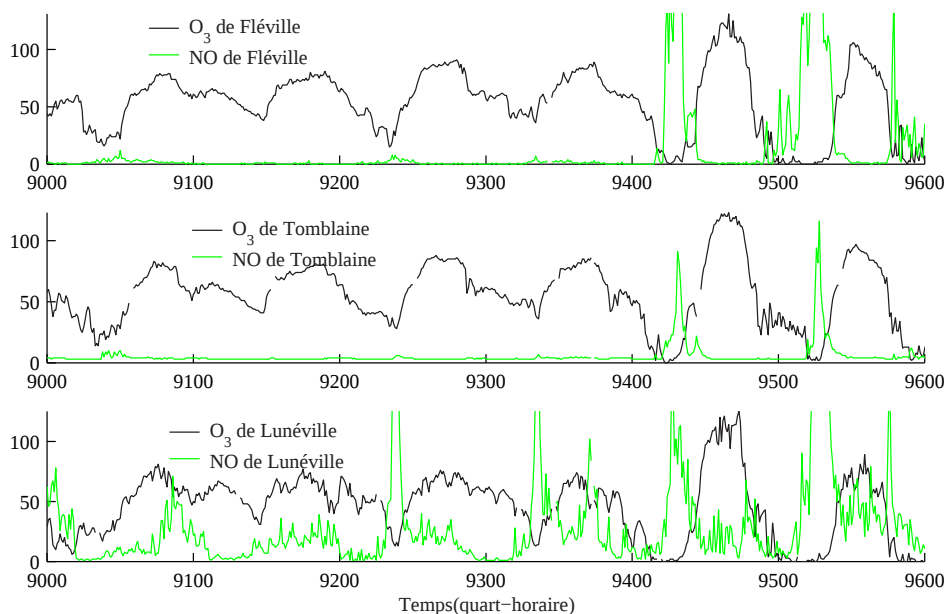


Figure 5.6 – Concentration de  $NO$  des sites de Brabois, Dan et Fléville sur une période de quelques jours

la concentration d'ozone d'un site et celles délivrées par d'autres capteurs d'ozone sur d'autres sites, fonctionne bien pour les valeurs maximales obtenues pendant l'après-midi mais pose des problèmes en début et fin de journée. On a constaté également que le fait d'ajouter d'autres

Figure 5.7 – Concentration d’ozone et de  $NO_2$  des sites de Fléville, Tomblaine et Lunéville.Figure 5.8 – Concentration d’ozone et de  $NO$  des sites de Fléville, Tomblaine et Lunéville.

variables explicatives dans le modèle telles que les concentrations en  $NO_2$  et  $NO$  permet de pallier en partie aux limitations du modèle précédent et l’erreur était inférieure à  $20\mu g/m^3$  tout en sachant que l’erreur relative de mesure des capteurs d’ozone est de 15%. Cependant, on doit chercher un ensemble de modèles pour les différents capteurs. En revanche, pour les systèmes avec un grand nombre de capteurs, établir de tels modèles pour les variables paraît moins immédiat.

De plus, le choix de modèle pour la génération de résidus est arbitraire et il se peut très bien que les corrélations entre certaines variables ne soient pas prise en compte. Pour cette raison, on utilise l'ACP car avec cette méthode toutes les corrélations entre les variables sont prisent en compte.

Ainsi, l'application de l'analyse en composantes principales devrait nous permettre de mieux prendre en compte les changement dans les la nature de la dépendance entre ozone et NO<sub>2</sub> et NO.

## 5.5 Application de l'ACP linéaire

Dans cette section, nous allons présenter l'application de l'ACP à la détection et la localisation de défauts de capteurs mesurant l'ozone et les oxydes d'azote. Dans cette application 18 variables ont été considérées. Les sites de mesure choisis sont géographiquement voisins et sont placés dans des sites ruraux, périurbains et urbains (fig 5.2).

La matrice  $X$  contient donc 18 variables,  $v_1, \dots, v_{18}$  correspondant, respectivement, à l'ozone O<sub>3</sub>, le monoxyde d'azote NO et le dioxyde d'azote NO<sub>2</sub> pour six stations de mesure.

Une première analyse des corrélations a tout d'abord été effectuée graphiquement et le résultat est reporté sur la figure (fig 5.9). Sur cette figure les capteurs ont été ordonnés tels que ceux qui présentent les plus fortes inter-corrélations soient les plus proches les uns des autres formant ainsi des groupes. Les plus fortes corrélations sont ainsi distribuées le long de la diagonale. Ces résultats indiquent par exemple que les capteurs 1, 4, 7, 10, 13 et 16 forment un groupe de capteurs avec des coefficients d'inter-corrélations supérieurs à 85%. Ces variables représentent les concentrations d'ozones des différents sites.

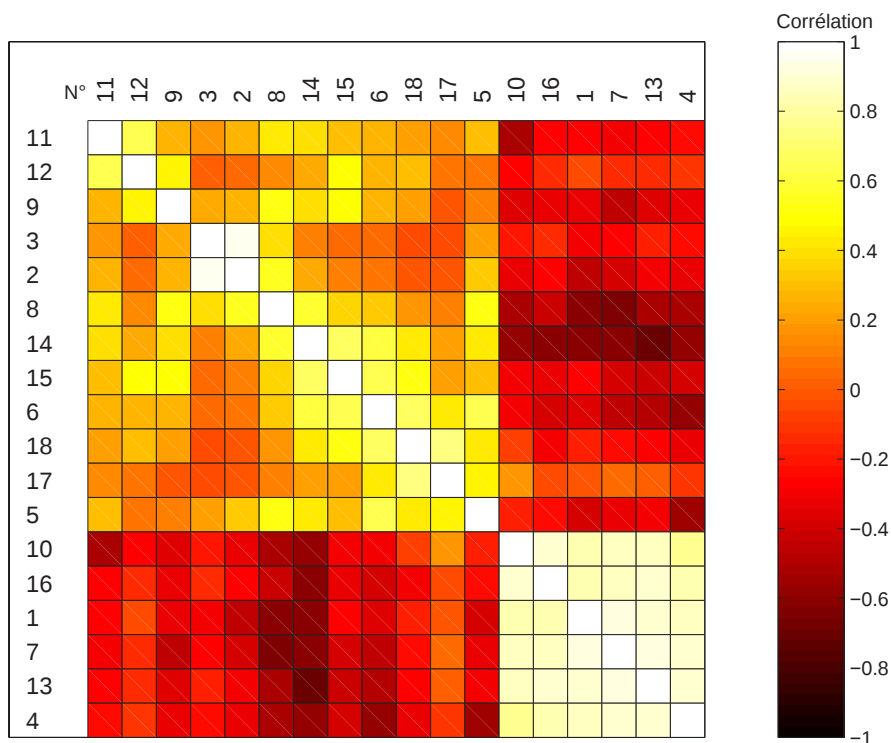


Figure 5.9 – Degrés de corrélation entre les capteurs (réseau AIRLOR).

Pour l'identification du modèle ACP, la matrice des corrélations des mesures est calculée. Sur le tableau (tab 5.2), on présente les valeurs propres de la matrice des corrélations des données ainsi que les parts de corrélation expliquée par chaque valeur propre. Pour définir la structure du modèle ACP on doit déterminer le nombre de composantes à retenir dans ce modèle en utilisant le critère de la variance non reconstruite. Le tableau (tab 5.3) représente les variances non reconstruites des différentes variables pour différentes valeurs de  $\ell$  et la figure (fig 5.10) présente l'évolution du critère de la variance non reconstruite en fonction du nombre de composantes. Ainsi, pour notre application, un modèle ACP à sept composantes a été retenu ce qui explique 91% des corrélations entre les variables. Alors qu'en utilisant le critère sur les valeurs propres on trouve un  $\ell = 5$  en expliquant uniquement 85% des corrélations. Il faut rappeler que ce dernier critère a donné le même nombre  $\ell$  que le critère de la variance non reconstruite appliqué sur les exemples de simulation dans le chapitre 2. Ainsi, en pratique on montre bien les limites d'un tel critère.

| PC | Valeur propre | Corrélation expliquée (%) | Cumul (%) |
|----|---------------|---------------------------|-----------|
| 1  | 7.29          | 40.52                     | 40.52     |
| 2  | 3.30          | 18.35                     | 58.88     |
| 3  | 1.81          | 10.09                     | 68.97     |
| 4  | 1.26          | 7.02                      | 76.00     |
| 5  | 1.21          | 6.75                      | 82.75     |
| 6  | 0.93          | 5.17                      | 87.93     |
| 7  | 0.64          | 3.56                      | 91.50     |
| 8  | 0.46          | 2.59                      | 94.09     |
| 9  | 0.26          | 1.47                      | 95.57     |
| 10 | 0.22          | 1.25                      | 96.82     |
| 11 | 0.17          | 0.96                      | 97.78     |
| 12 | 0.16          | 0.90                      | 98.69     |
| 13 | 0.08          | 0.46                      | 99.16     |
| 14 | 0.06          | 0.37                      | 99.54     |
| 15 | 0.03          | 0.17                      | 99.71     |
| 16 | 0.02          | 0.16                      | 99.97     |
| 17 | 0.01          | 0.09                      | 99.97     |
| 18 | 0.00          | 0.03                      | 100.00    |

Table 5.2 – Valeurs propres de la matrice de corrélation

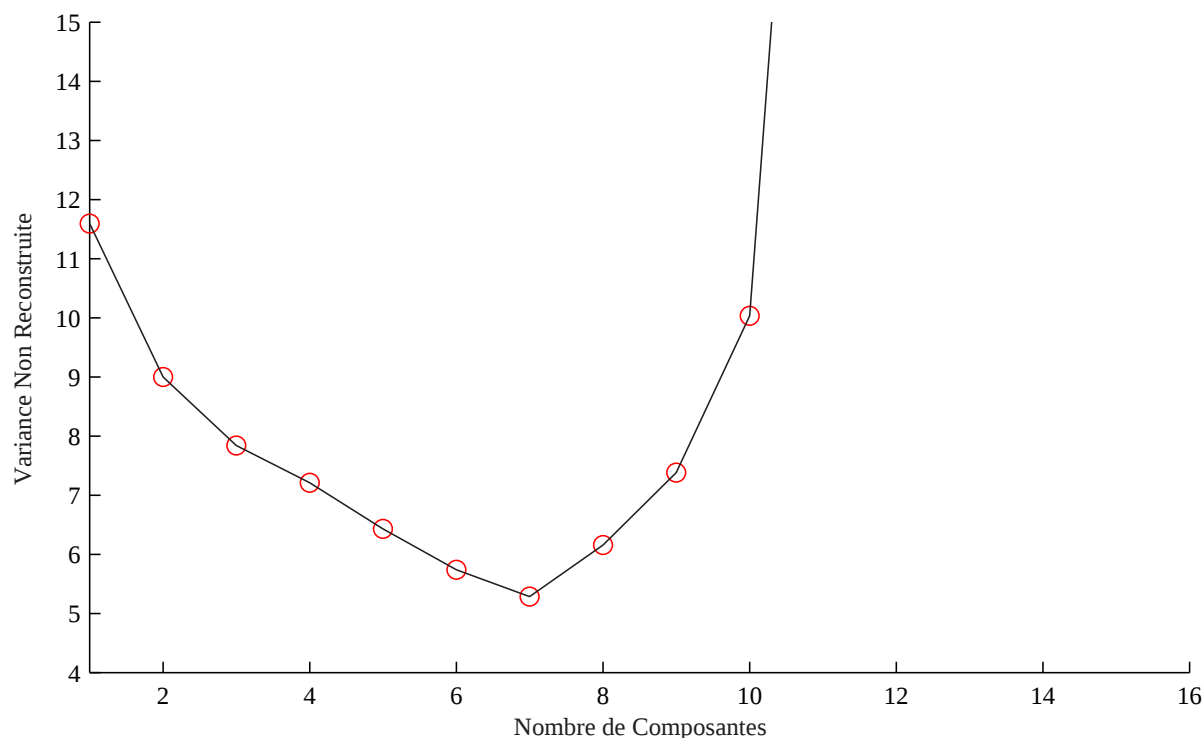
A partir du tableau (tab 5.3) et pour  $\ell = 7$ , on constate que toutes les variables peuvent être reconstruites mais il y a quelques difficultés envisageables pour les variables 9, 11, 12 et 15 qui ont des coefficients un peu élevés par rapport aux autres variables. Il faut noter que ces variables représentent, respectivement, le NO mesuré sur le site de Fléville, NO<sub>2</sub> et NO de Dan, et le NO<sub>2</sub> de St-Nicolas.

Ainsi, les deux matrices (5.1) et (5.2) représentent respectivement la matrice des premiers vecteurs propres retenus dans le modèle ACP et la matrice des derniers vecteurs propres représentant l'espace résiduel.

| $i$                         | $\ell$ |      |      |      |             |      |      | $\xi_i^T \Sigma \xi_i$ | $\rho_i(7) < \xi_i^T \Sigma \xi_i$ |
|-----------------------------|--------|------|------|------|-------------|------|------|------------------------|------------------------------------|
|                             | 3      | 4    | 5    | 6    | 7           | 8    | 9    |                        |                                    |
| 1                           | 0.12   | 0.12 | 0.09 | 0.09 | <b>0.09</b> | 0.08 | 0.08 | 1                      | ✓                                  |
| 2                           | 0.29   | 0.16 | 0.12 | 0.11 | <b>0.11</b> | 0.01 | 0.01 | 1                      | ✓                                  |
| 3                           | 0.46   | 0.27 | 0.17 | 0.15 | <b>0.15</b> | 0.04 | 0.04 | 1                      | ✓                                  |
| 4                           | 0.21   | 0.21 | 0.17 | 0.17 | <b>0.10</b> | 0.10 | 0.10 | 1                      | ✓                                  |
| 5                           | 0.55   | 0.55 | 0.50 | 0.50 | <b>0.37</b> | 0.37 | 0.38 | 1                      | ✓                                  |
| 6                           | 0.51   | 0.50 | 0.50 | 0.49 | <b>0.41</b> | 0.41 | 0.44 | 1                      | ✓                                  |
| 7                           | 0.05   | 0.05 | 0.04 | 0.04 | <b>0.04</b> | 0.03 | 0.03 | 1                      | ✓                                  |
| 8                           | 0.44   | 0.39 | 0.32 | 0.30 | <b>0.29</b> | 0.26 | 0.26 | 1                      | ✓                                  |
| 9                           | 0.63   | 0.62 | 0.59 | 0.59 | <b>0.59</b> | 1.72 | 1.87 | 1                      | ✓                                  |
| 10                          | 0.05   | 0.05 | 0.05 | 0.04 | <b>0.04</b> | 0.04 | 0.04 | 1                      | ✓                                  |
| 11                          | 0.81   | 0.82 | 0.76 | 0.54 | <b>0.54</b> | 0.54 | 0.78 | 1                      | ✓                                  |
| 12                          | 0.83   | 0.85 | 0.60 | 0.58 | <b>0.57</b> | 0.57 | 0.94 | 1                      | ✓                                  |
| 13                          | 0.05   | 0.05 | 0.04 | 0.04 | <b>0.04</b> | 0.03 | 0.03 | 1                      | ✓                                  |
| 14                          | 0.54   | 0.50 | 0.43 | 0.38 | <b>0.37</b> | 0.37 | 0.49 | 1                      | ✓                                  |
| 15                          | 0.66   | 0.66 | 0.66 | 0.64 | <b>0.61</b> | 0.61 | 0.89 | 1                      | ✓                                  |
| 16                          | 0.16   | 0.15 | 0.13 | 0.10 | <b>0.09</b> | 0.09 | 0.08 | 1                      | ✓                                  |
| 17                          | 0.70   | 0.69 | 0.68 | 0.50 | <b>0.42</b> | 0.42 | 0.43 | 1                      | ✓                                  |
| 18                          | 0.68   | 0.50 | 0.50 | 0.40 | <b>0.37</b> | 0.37 | 0.39 | 1                      | ✓                                  |
| $\sum_{i=1}^m \rho_i(\ell)$ | 7.84   | 7.21 | 6.43 | 5.73 | <b>5.28</b> | 6.15 | 7.38 |                        |                                    |

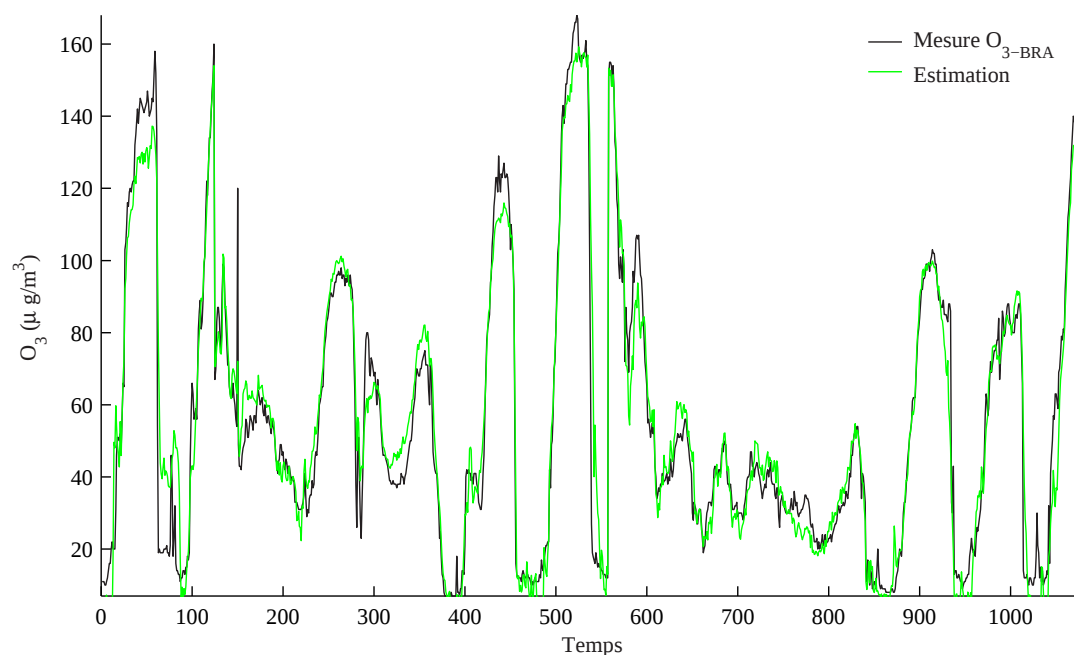
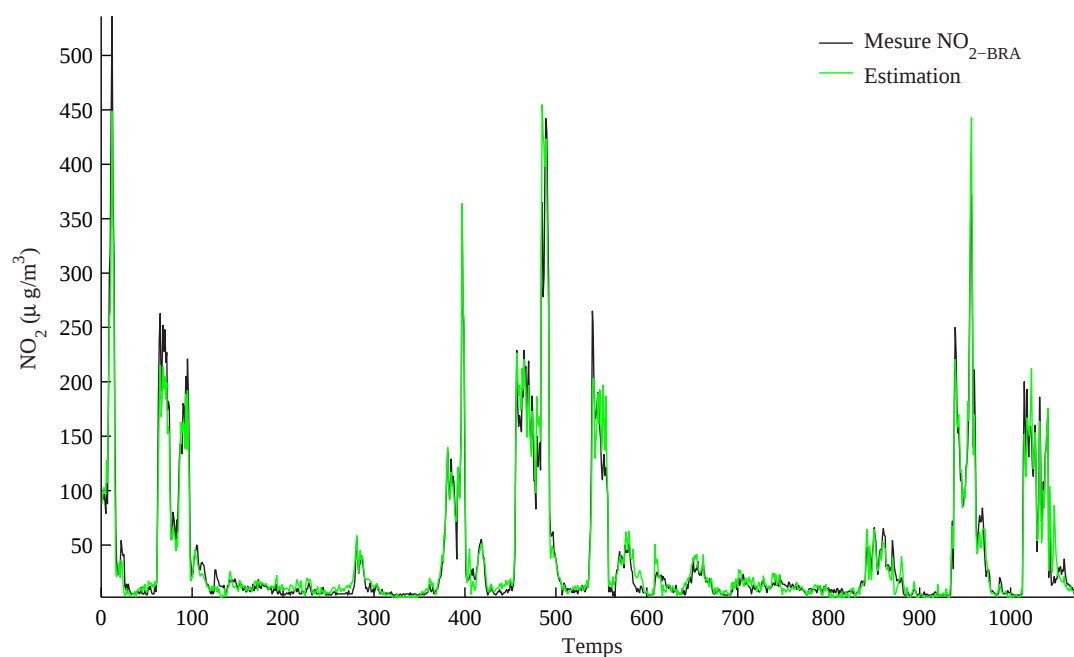
Table 5.3 – Variances des erreurs de reconstruction des différentes variables

$$\hat{P} = \begin{pmatrix} 0.32 & -0.20 & 0.10 & -0.01 & -0.13 & -0.06 & 0.03 \\ -0.21 & -0.21 & -0.43 & -0.25 & -0.12 & -0.05 & -0.00 \\ -0.20 & -0.18 & -0.41 & -0.34 & -0.19 & -0.09 & 0.02 \\ 0.31 & -0.17 & -0.13 & 0.01 & -0.16 & -0.07 & -0.30 \\ -0.07 & -0.41 & -0.09 & 0.10 & 0.27 & 0.02 & 0.55 \\ -0.19 & -0.16 & 0.42 & -0.07 & 0.01 & -0.20 & 0.48 \\ 0.33 & -0.19 & 0.00 & -0.05 & -0.06 & -0.04 & 0.03 \\ -0.23 & -0.20 & -0.31 & 0.22 & 0.23 & 0.15 & -0.08 \\ -0.22 & -0.17 & -0.04 & -0.21 & -0.33 & -0.19 & -0.04 \\ 0.33 & -0.21 & -0.03 & -0.04 & 0.02 & -0.06 & 0.04 \\ -0.15 & -0.16 & -0.08 & 0.56 & -0.28 & 0.39 & -0.02 \\ -0.15 & -0.10 & 0.23 & 0.21 & -0.65 & 0.08 & 0.08 \\ 0.31 & -0.24 & -0.10 & 0.04 & -0.04 & 0.04 & 0.06 \\ -0.20 & -0.29 & 0.07 & 0.26 & 0.27 & -0.26 & -0.18 \\ -0.13 & -0.30 & 0.22 & 0.16 & 0.07 & -0.47 & -0.45 \\ 0.32 & 0.16 & -0.06 & 0.10 & -0.11 & -0.14 & 0.15 \\ 0.04 & -0.36 & 0.15 & -0.19 & 0.21 & 0.56 & -0.22 \\ -0.12 & -0.23 & 0.40 & -0.42 & -0.04 & 0.26 & -0.15 \end{pmatrix} \quad (5.1)$$

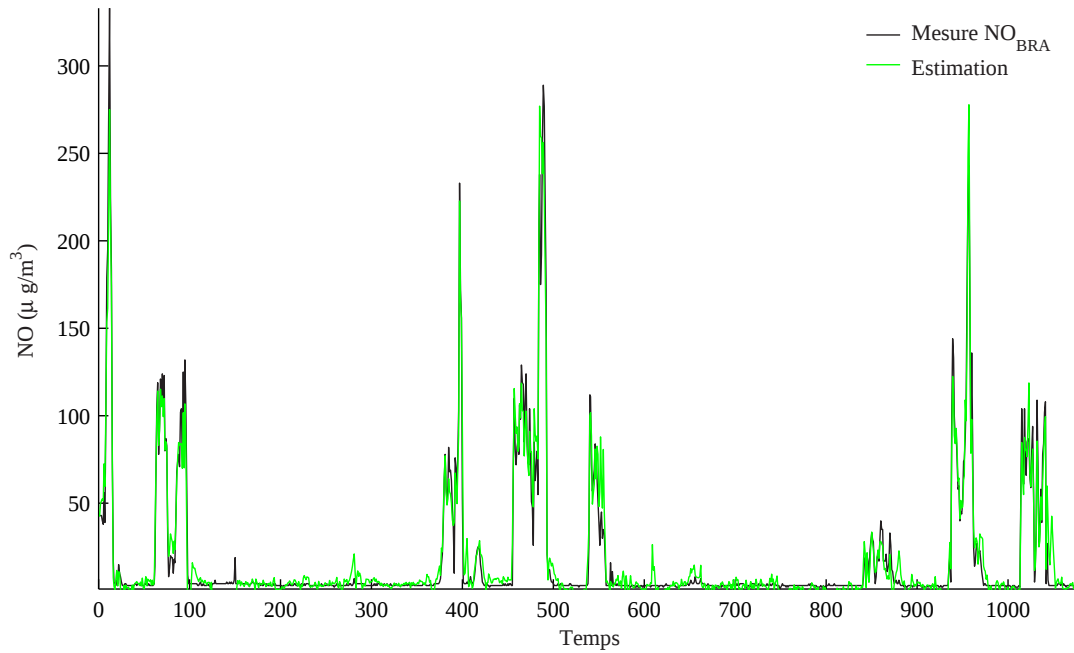
Figure 5.10 – Evolution de la variance non reconstruite en fonction de  $\ell$ .

$$\tilde{P} = \begin{pmatrix} -0.09 & 0.11 & -0.08 & 0.01 & 0.04 & 0.48 & -0.68 & 0.00 & 0.18 & -0.01 & -0.22 \\ 0.30 & 0.00 & 0.04 & 0.04 & 0.01 & -0.05 & 0.08 & -0.09 & 0.01 & -0.05 & -0.71 \\ 0.62 & 0.30 & 0.07 & 0.04 & 0.11 & -0.01 & 0.04 & -0.26 & 0.07 & -0.02 & 0.62 \\ 0.03 & 0.10 & 0.04 & -0.01 & 0.30 & 0.40 & 0.53 & 0.35 & 0.13 & 0.13 & 0.07 \\ -0.04 & -0.27 & -0.19 & -0.18 & -0.34 & 0.25 & 0.18 & 0.18 & 0.05 & 0.08 & 0.04 \\ 0.07 & 0.22 & 0.37 & 0.33 & 0.38 & 0.04 & 0.10 & -0.02 & -0.00 & 0.01 & -0.01 \\ 0.08 & 0.05 & -0.02 & 0.07 & -0.13 & -0.14 & 0.11 & -0.61 & 0.14 & 0.59 & 0.08 \\ -0.32 & -0.15 & 0.26 & -0.26 & 0.53 & -0.05 & -0.19 & -0.12 & 0.05 & 0.23 & 0.03 \\ -0.78 & 0.13 & -0.13 & 0.13 & -0.16 & -0.06 & 0.11 & -0.02 & -0.01 & -0.01 & -0.00 \\ -0.02 & 0.04 & -0.05 & -0.01 & 0.09 & -0.20 & -0.10 & 0.21 & -0.82 & 0.15 & -0.06 \\ 0.06 & 0.41 & 0.27 & 0.09 & -0.32 & 0.07 & -0.00 & 0.01 & -0.14 & 0.03 & -0.00 \\ 0.15 & -0.43 & -0.33 & -0.09 & 0.26 & -0.10 & -0.00 & -0.01 & -0.01 & 0.03 & 0.02 \\ -0.05 & -0.07 & 0.12 & -0.09 & 0.08 & 0.03 & 0.16 & -0.46 & -0.10 & -0.70 & 0.13 \\ 0.15 & 0.44 & -0.57 & -0.05 & 0.15 & -0.14 & -0.00 & -0.08 & 0.06 & -0.12 & 0.04 \\ 0.09 & -0.38 & 0.34 & 0.06 & -0.28 & 0.03 & -0.07 & -0.00 & -0.06 & -0.00 & 0.00 \\ -0.01 & 0.07 & 0.19 & -0.13 & -0.04 & -0.62 & -0.10 & 0.37 & 0.42 & -0.09 & -0.01 \\ -0.02 & -0.17 & -0.11 & 0.52 & 0.03 & -0.15 & -0.02 & 0.12 & 0.12 & -0.06 & 0.00 \\ 0.07 & 0.20 & 0.13 & -0.64 & -0.08 & -0.01 & 0.03 & 0.00 & -0.01 & 0.01 & 0.01 \end{pmatrix} \quad (5.2)$$

Les trois figures (fig 5.11), (fig 5.12) et (fig 5.13) présentent, respectivement, les mesures et les estimations d'ozone, de  $\text{NO}_2$  et de  $\text{NO}$  pour une station de mesure. Les estimations sont données par le modèle ACP.

Figure 5.11 – Mesure et estimation de l'ozone de Brabois ( $v_1$ ) par le modèle ACP linéaire.Figure 5.12 – Mesure et estimation du  $\text{NO}_2$  de Brabois ( $v_2$ ).

En tenant compte de la nature du processus modélisé, les résultats obtenus sont très satisfaisants avec le modèle ACP obtenu. En effet, on estime la plupart des pics de  $\text{NO}$ ,  $\text{O}_3$  et  $\text{NO}_2$  qui sont prépondérants pour les procédures d'alerte. De plus, dans le cas des oxydes d'azote ( $\text{NO}$  et

Figure 5.13 – Mesure et estimation du NO de Brabois ( $v_3$ ).

$\text{NO}_2$ ) qui sont des polluants plus localisés, et plus difficile à modéliser, l'estimation de ces deux grandeurs reste correcte pour les faibles valeurs ainsi que pour les valeurs élevées.

Basés sur le modèle ACP déjà obtenu, les indices de détection et de localisation de capteurs défaillants peuvent être calculés en ligne. Ainsi, la figure (fig 5.14) présente l'évolution du  $SPE$  et  $SPE$  filtré en présence d'un défaut affectant la variable  $v_7$  ( $\text{O}_3$  de Fléville). L'amplitude du défaut simulé s'élève à environ 20% de la plage de variation de la mesure  $v_7$  entre les instants 800 et 1080. Comme on peut le constater sur la figure (fig 5.14), le défaut n'est pas détecté sur  $SPE$  ni sur  $SPE$  filtré à cause des erreurs de modélisation. Face à ce problème, trois solutions sont envisageables :

- la première solution consiste à considérer plus de composantes dans le modèle afin de réduire les erreurs de modélisation, mais en adoptant cette solution certains défauts ne seront plus détectables car ils seront directement projetés dans le sous-espace principal,
- la deuxième solution consiste à calculer des indices de détection dans les différents sous-espaces résiduels (somme successive des carrés des dernières composantes principales),
- la troisième solution utilise une approche non-linéaire pour la modélisation du processus.

Nous avons opté pour la solution utilisant l'indice de détection  $D_i$  que nous avons développé dans le troisième chapitre de cette thèse. En appliquant la procédure de détection utilisant la somme des carrés des dernières composantes principales, le défaut a été très nettement détecté sur l'indice  $D_2$  filtré (somme des carrés des deux dernières composantes) (fig 5.15).

Une fois le défaut détecté, nous cherchons à localiser la variable en défaut. Pour cela, deux approches ont été appliquées : la première consiste à exploiter la procédure de reconstruction et la deuxième approche consiste à calculer les contributions des différentes variables à l'indice de détection  $D_2$ .

Ainsi, le résultat de l'application de la première approche est illustré sur les deux figures (fig 5.16) et (fig 5.17).

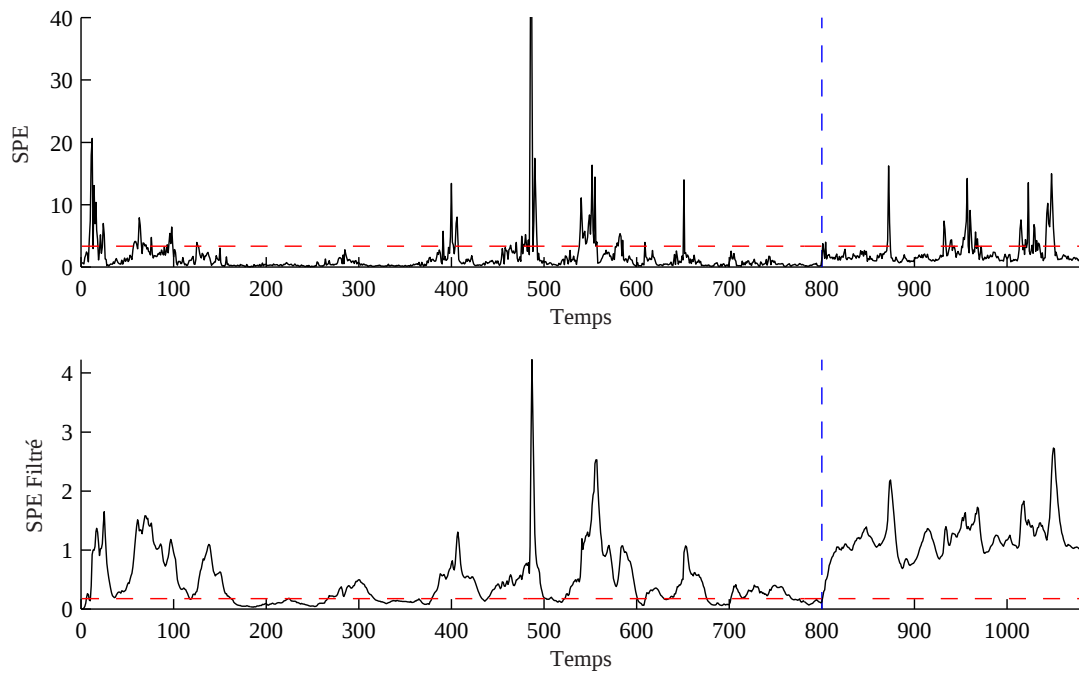


Figure 5.14 – Evolution de l'erreur quadratique en présence d'un défaut affectant la variable 7.

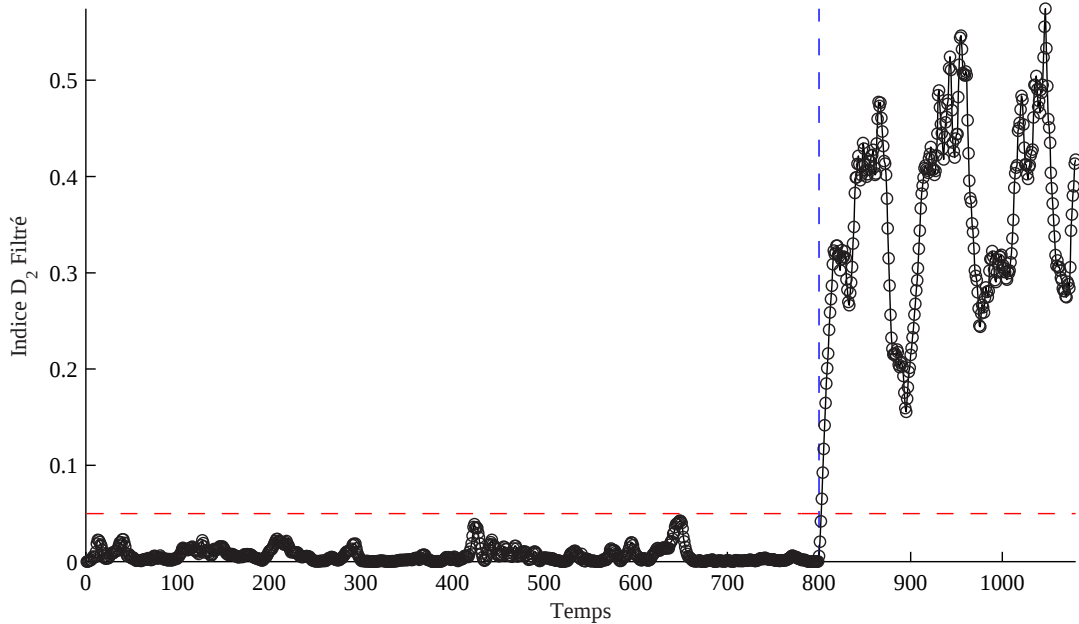


Figure 5.15 – Evolution de l'indice  $D_2$  en présence d'un défaut affectant la variable  $v_7$ .

La figure (fig 5.16) présente l'évolution des différents indices de détection calculés après la reconstruction d'une variable parmi l'ensemble des variables à surveiller. Il est clair que l'indice calculé après la reconstruction de la variable  $v_7$  ne présente pas un dépassement de son seuil ce

qui implique que la variable considérée est la variable en défaut. Ce résultat est confirmé par la figure (fig 5.17).

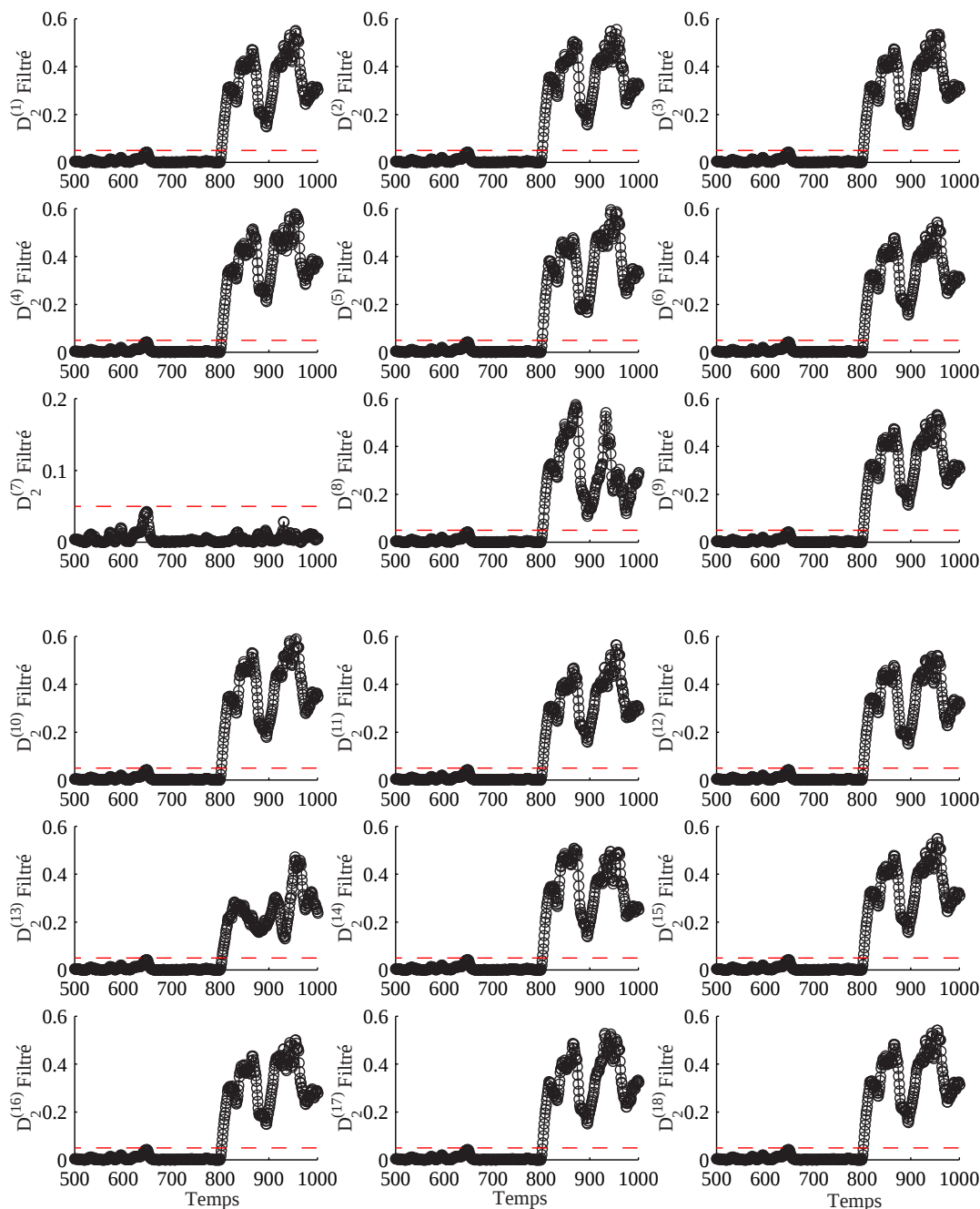


Figure 5.16 – Localisation de la variable  $v_7$  par reconstruction des 18 variables.

Le résultat de l'application de la deuxième approche utilisant le calcul des contribution des différentes variables à l'indice de détection est illustré sur la figure (fig 5.18). La variable  $v_7$  est la variable qui présente la plus forte contribution à l'indice  $D_2$ , ce qui indique que c'est la variable incriminée.

Après la localisation de la variable en défaut, cette dernière est reconstruite en utilisant le

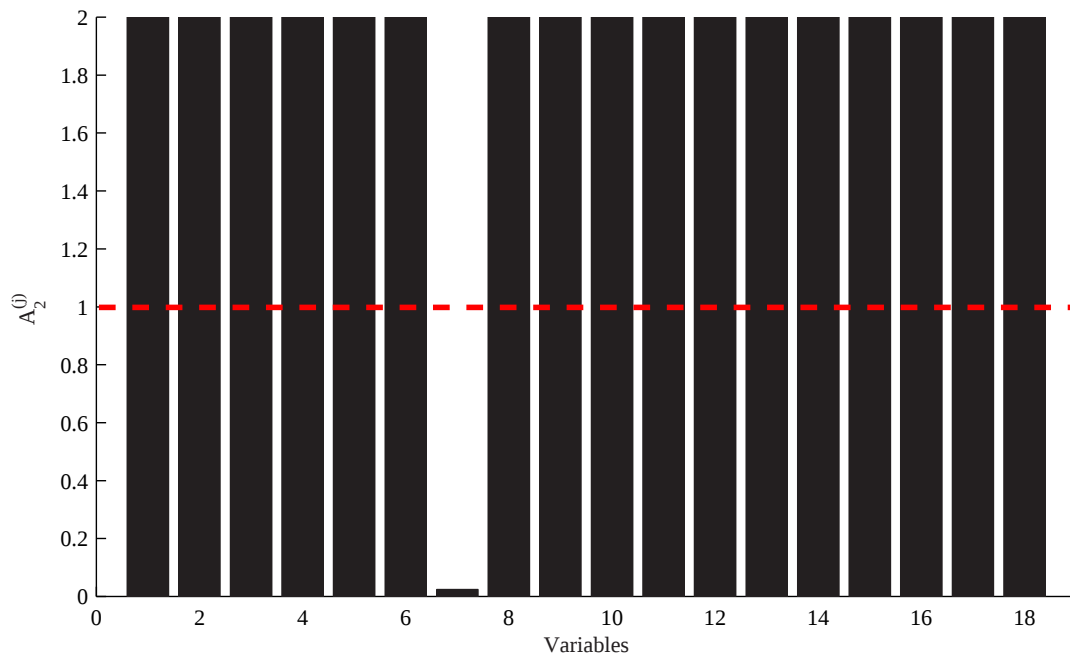


Figure 5.17 – Localisation par l'indice  $A_2^{(j)}$  calculé après la reconstruction de la  $j^{eme}$  variable.

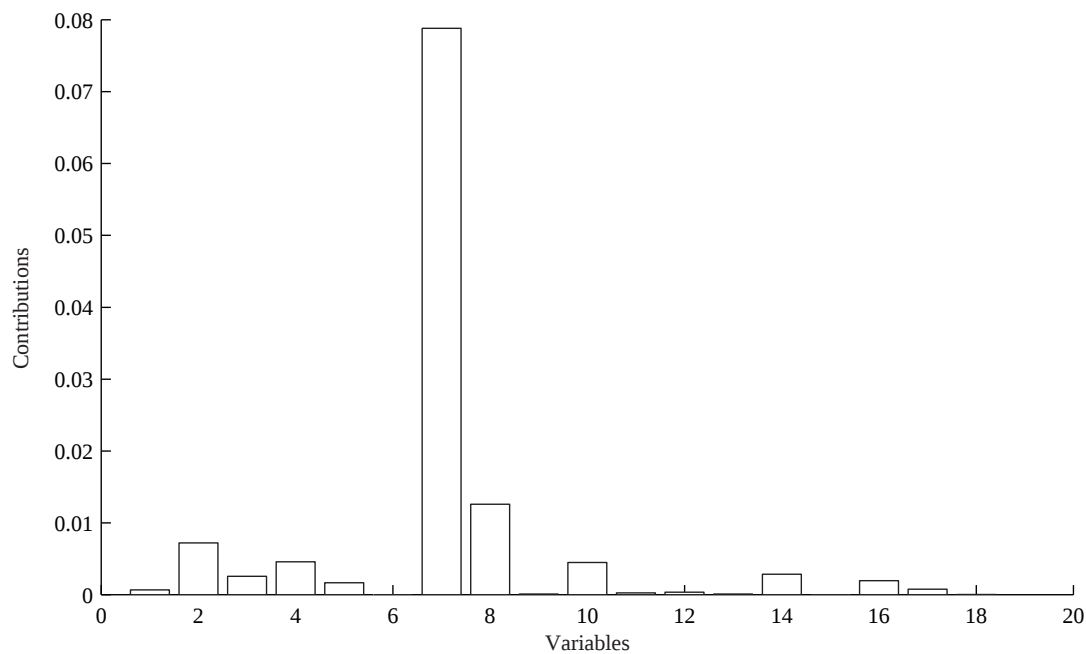


Figure 5.18 – Localisation par calcul des contributions des variables à l'indice  $D_2$ .

modèle ACP et les mesures des autres capteurs. La figure (fig 5.19) représente l'évolution de la variable  $v_7$  représentant l'ozone mesuré par la station de Fléville, l'évolution de cette variable avec le défaut simulé et la reconstruction de cette dernière.

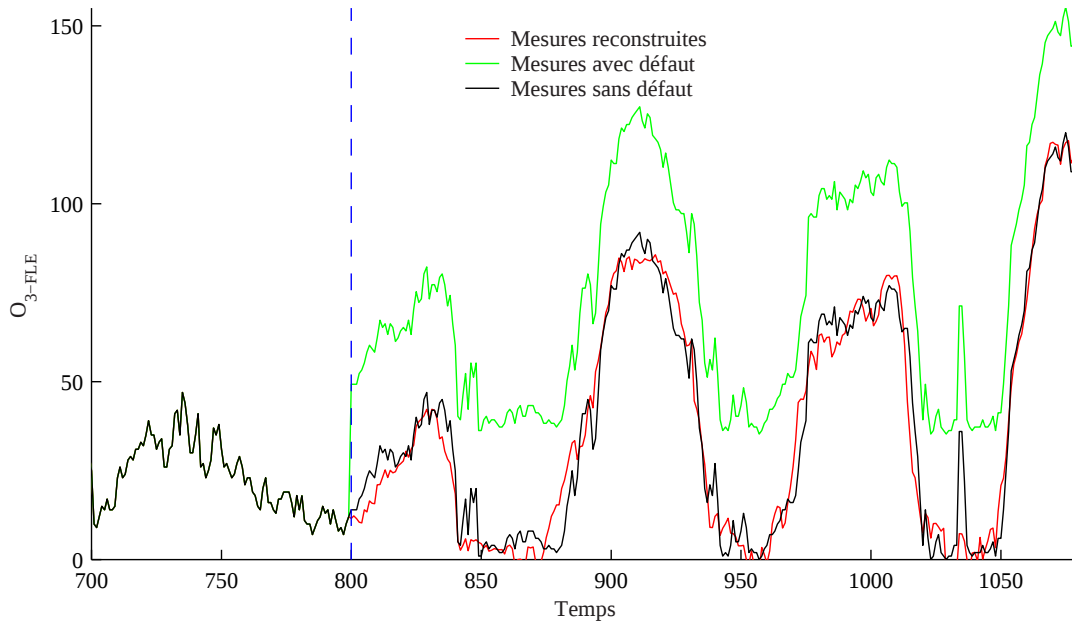


Figure 5.19 – Reconstruction de la variable en défaut cas de l'ozone de Fléville  $O_3-FLE$ .

Un autre défaut a été simulé sur une autre variable  $v_{10}$  représentant l'ozone de St-Nicolas. Le défaut simulé est un biais avec une amplitude d'environ 20% de la plage de variation de cette variable entre les instant 800 et 1080.

L'application de la procédure de détection a permis de détecter le défaut à l'instant 805 sur l'indice  $D_3$  filtré qui est représenté sur la figure (fig 5.20). Il faut noter que le retard à la détection est assez faible de plus que la détection est nette.

Pour la localisation de la variable incriminée, les figures (fig 5.21) et (fig 5.22) présentent le résultat de l'application de l'approche basée sur le principe de reconstruction. Ainsi, sur la figure (fig 5.21) l'évolution des différents indices indique que c'est l'indice  $D_3^{(10)}$  qui ne présente pas de dépassement de son seuil ce qui indique que la variable  $v_{10}$  est la variable incriminée comme le montre la figure (fig 5.22).

La figure (fig 5.23) présente les contribution des différentes variables à l'indice de détection  $D_3$  filtré calculé à l'instant de détection et montre que  $v_{10}$  est la variable en défaut.

On a montré que la procédure de détection et localisation marche très bien pour les variables  $v_7$  et  $v_{10}$ , vue les corrélations entre les différents capteurs d'ozone, on est à peu près sûr que ça marche aussi pour tous les  $O_3$  surtout avec les résultats du tableau (tab 5.3) et la matrice des derniers vecteurs propres (5.2).

Maintenant intéressons nous aux  $NO_x$  qui sont des signaux plus difficiles à modéliser.

Dans la suite nous allons considérer le cas des oxydes d'azote. Dans un premier temps, nous allons simuler un défaut affectant la variable  $v_2$  représentant le  $NO_2$  de Brabois avec une amplitude d'environ 20% de la plage de variation de cette variable entre les instants 800 et 1080. A partir de la matrice des derniers vecteurs propres il est évident que le défaut sera détecté en utilisant l'indice  $D_1$  car cette variable apparaît avec un coefficient de  $-0.7$  sur le dernier vecteur propre. Ainsi, l'évolution de l'indice  $D_1$  filtré en présence de ce défaut est illustré sur la figure (fig 5.25).

Comme dans le cas de l'ozone, pour la localisation on applique les deux approches de recons-

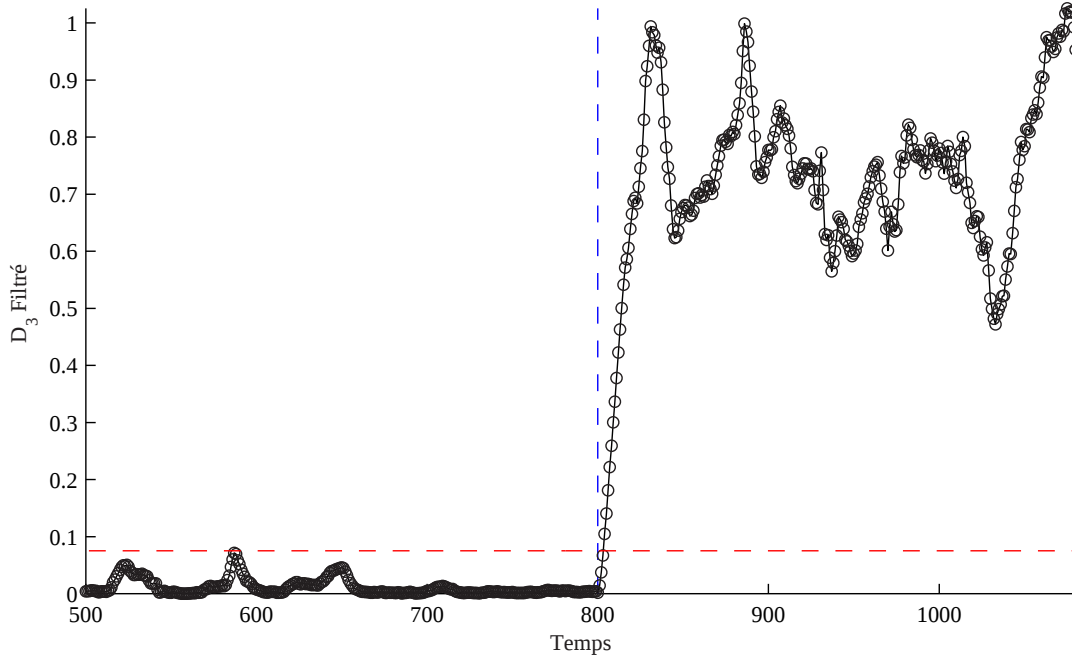


Figure 5.20 – Evolution de l'indice  $D_3$  filtré en présence d'un défaut affectant la variable  $v_{10}$ .

truction et des contributions à l'indice de détection  $D_1$ . Le résultat de l'application de l'approche de reconstruction à cet indice est représenté sur la figure (fig 5.26).

Ainsi, la variable reconstruite pour laquelle l'indice de détection ne présente pas un dépassement de son seuil est bien la variable  $v_2$ .

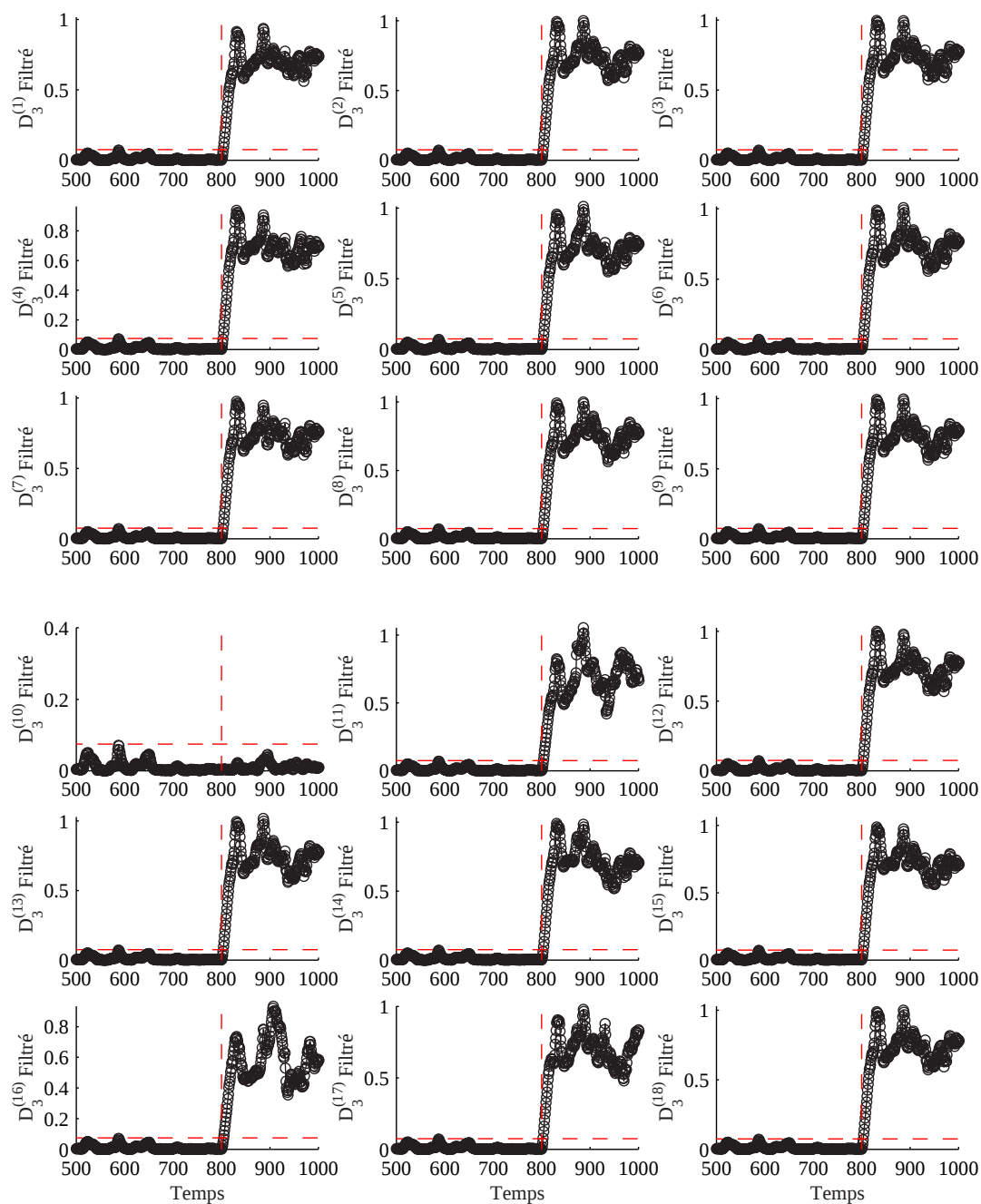
Le même résultat peut être obtenu en calculant l'indice  $A_1$  où la variable  $v_2$  est la variable pour laquelle l'indice  $A_1$  est inférieur à 1 (fig 5.27) ce qui indique bien que c'est la variable incriminée.

Concernant le calcul des contributions à l'indice  $D_1$ , il est représenté sur la figure (fig 5.28) où la variable  $v_2$  présente la plus forte contribution. Cependant  $v_3$  a une contribution assez forte.

Une fois la variable en défaut localisée, elle est reconstruite en utilisant le modèle ACP et les mesures des autres capteurs. Dans le cas du dioxyde d'azote du site de Brabois, la reconstruction de cette variable en défaut est représentée sur la figure (fig 5.29). En tenant compte de la nature de cette variable, le résultat obtenu est très satisfaisant.

Une dernière simulation pour le NO est présentée. Un défaut a été simulé sur la variable  $v_3$  représentant le NO de Brabois avec un défaut de même amplitude que précédemment et sur la même période. Ce défaut a été détecté sur l'indice  $D_1$  à l'instant 803 comme le montre la figure (fig 5.30). Pour la localisation de la variable en défaut, la procédure de reconstruction est appliquée. Le résultat est représenté sur la figure (fig 5.31). A partir de cette figure on constate que la localisation de cette variable pose un petit problème dans une certaine période. Sur l'indice  $D_1^{(2)}$ , par exemple aux instants 957, 958 et 959 et entre les instants 974 et 978 les valeurs de cet indice sont inférieures à son seuil de détection ce qui peut provoquer une ambiguïté entre les deux variables  $v_2$  et  $v_3$ . Pour résoudre ce problème nous devons tenir compte de la persistance du défaut. La reconstruction de la variable incriminée  $v$  (représentant le NO de Brabois) est illustrée sur la figure (fig 5.32).

Les résultats que nous venons de présenter ici avec ces défauts simulés sont très intéressants

Figure 5.21 – Localisation de la variable  $v_{10}$  par reconstruction des 18 variables.

du fait de la nature des polluants traités. En effet les oxydes d'azote sont des polluants localisés qui dépends de phénomènes non mesurés comme les COV par exemple et la circulation des voiture.

Toutefois, il faut noter que certaines variables seront plus difficilement reconstituables. Ces variables ont les coefficients les plus forts dans le tableau des variances non reconstituées.

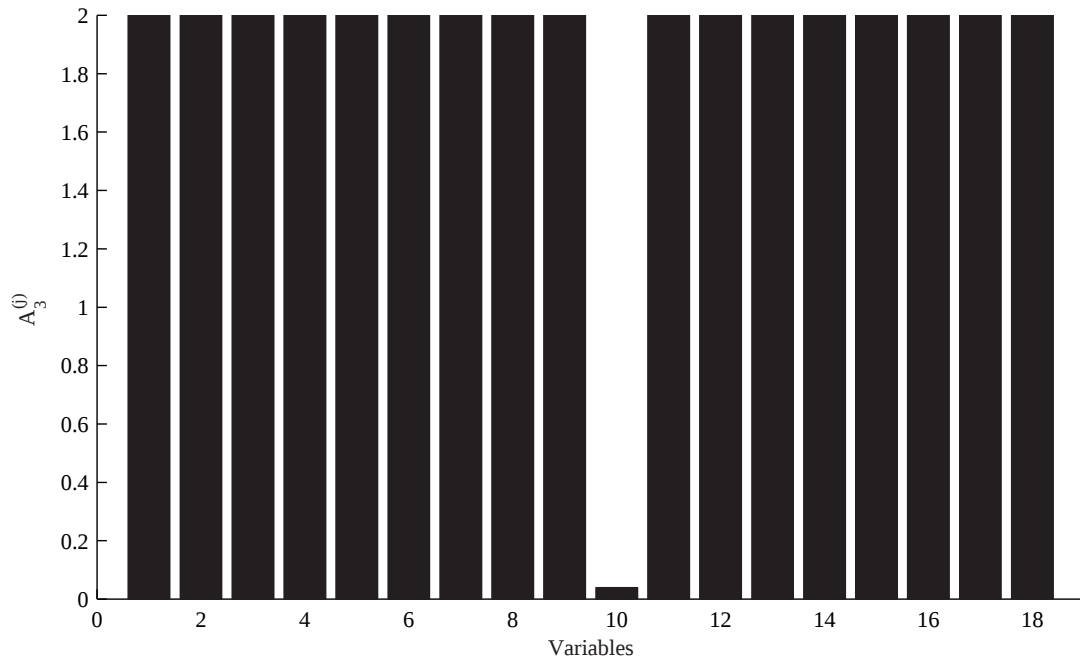


Figure 5.22 – Localisation par l'indice  $A_2^{(j)}$  calculé après la reconstruction de la  $j^{eme}$  variable.

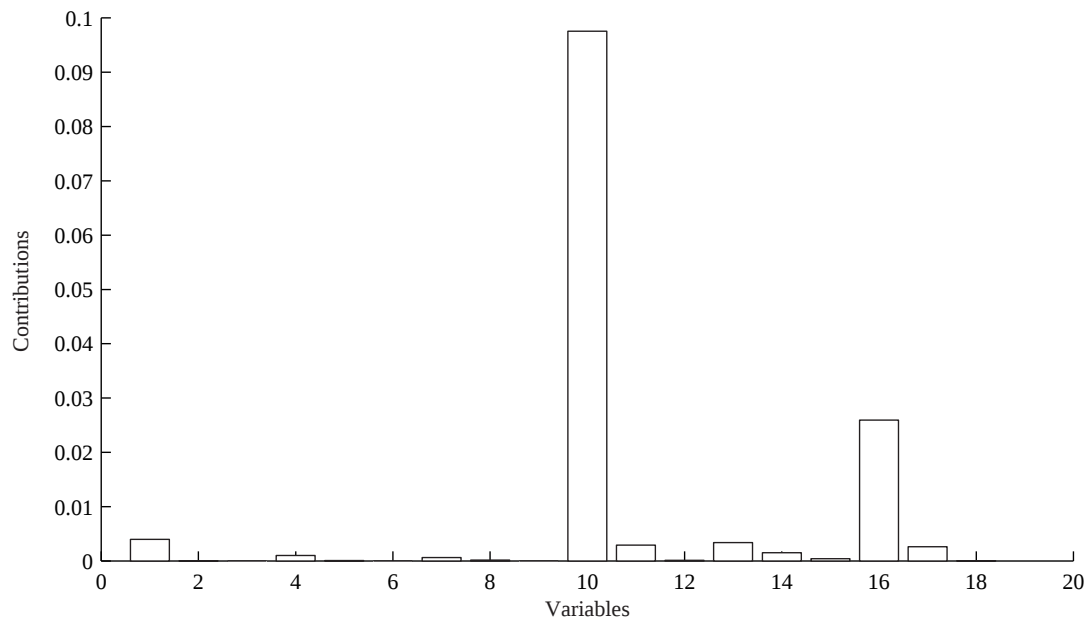
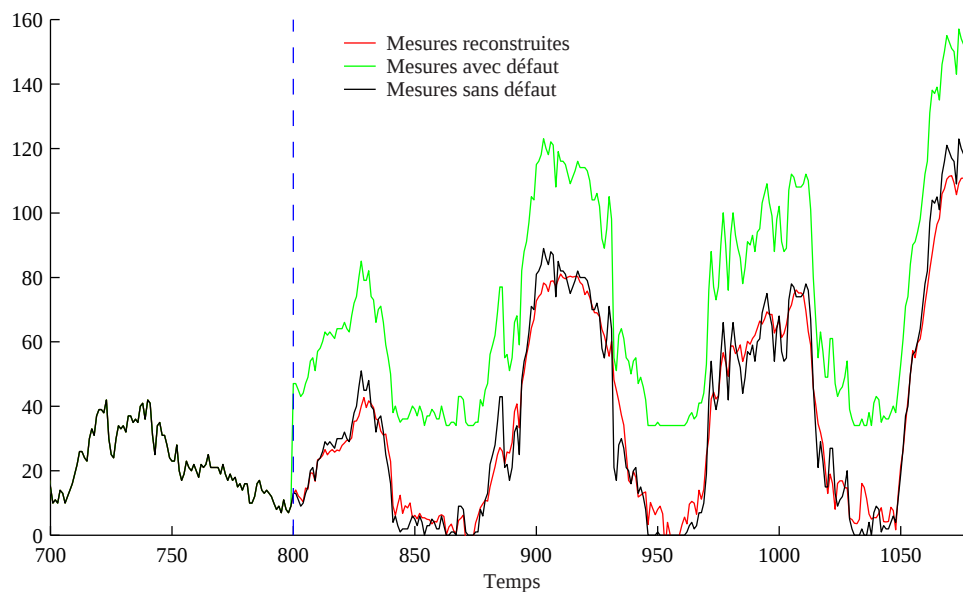
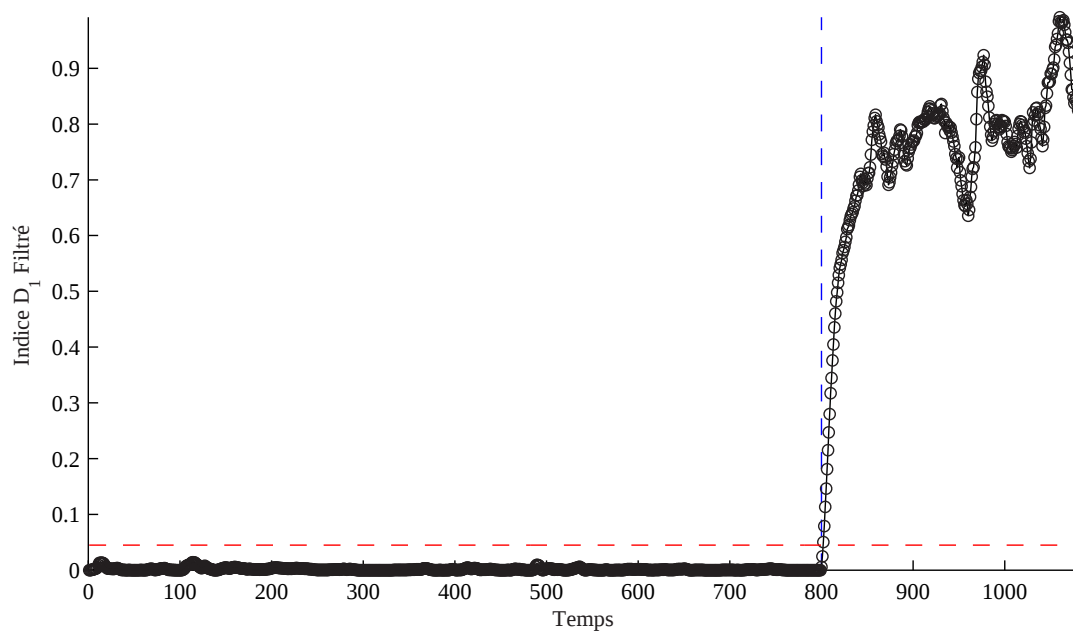


Figure 5.23 – Localisation de la variable  $v_{10}$  par calcul des contributions des variables à l'indice  $D_3$ .

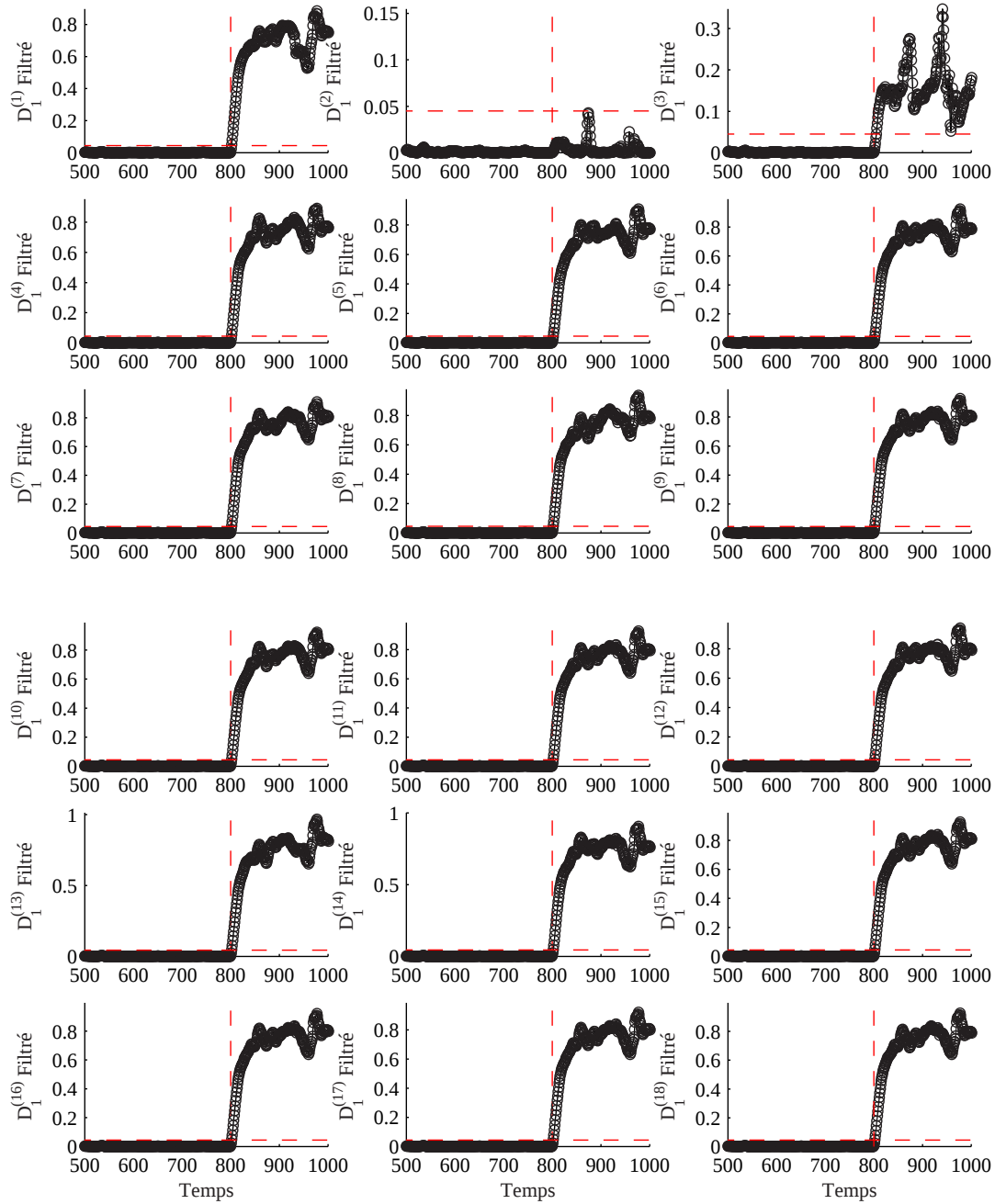
## 5.6 Conclusion

Ce chapitre a été consacré à l'application de l'analyse en composantes principales pour la détection et la localisation de défauts de capteurs d'un réseau de surveillance de la qualité de

Figure 5.24 – Reconstruction de la variable en défaut cas de l’ozone de St-Nicolas  $O_{3-STN}$ .Figure 5.25 – Evolution de l’indice  $D_1$  filtré en présence d’un défaut affectant la variable  $v_2$ .

l’air. Une partie du réseau a été prise en compte et une partie des données sur une période de plusieurs jours. Ainsi, 18 variables correspondant à six stations de mesures ont été considérées. Un modèle ACP a été identifié avec sept composantes en minimisant la variance de l’erreur de reconstruction.

Le modèle établi permet ainsi une prédiction satisfaisante des concentrations de polluant pour chacun des capteurs à partir les mesures fournies par l’ensemble des capteurs. On arrive à estimer

Figure 5.26 – Localisation de la variable  $v_2$  par reconstruction des 18 variables.

en plus des concentrations d’ozone, les concentrations en  $\text{NO}_2$  et  $\text{NO}$  qui sont des variables plus localisées que l’ozone et donc plus difficiles à modéliser.

Disposant d’un comportement prédit de chaque capteur, des tests de cohérence des mesures peuvent être appliqués afin de détecter, puis de localiser d’éventuelles anomalies de comportement des capteurs. Pour tester les approches de détection et de localisation proposées dans ce mémoire, nous avons simulé dans un premier temps, des défauts sur les variables représentant les concentrations d’ozone.

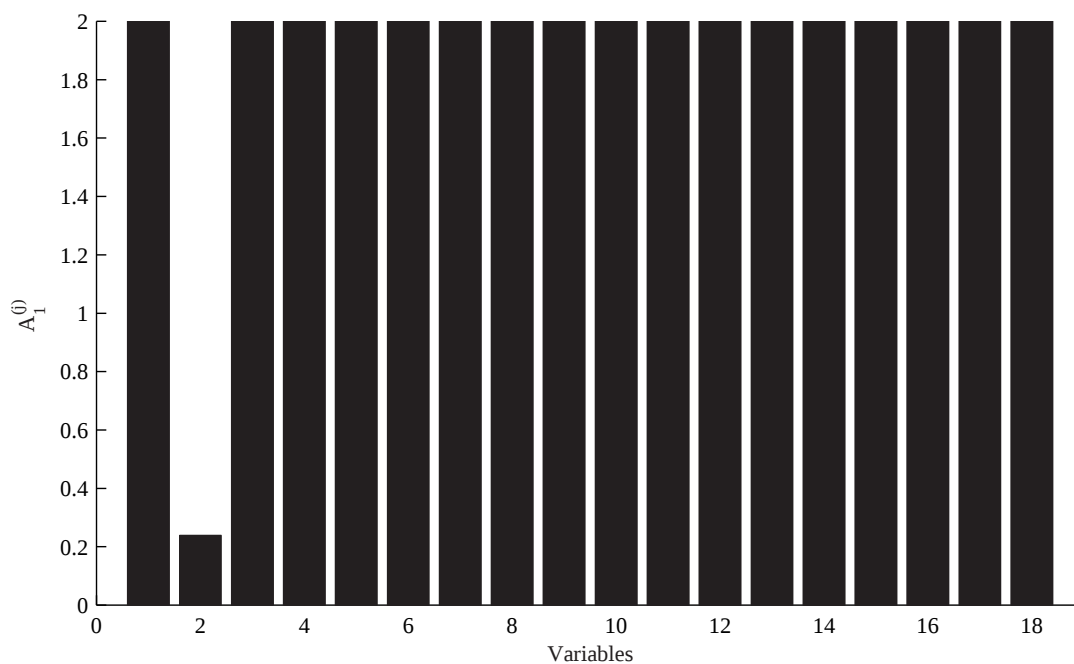


Figure 5.27 – Localisation par l'indice  $A_2^{(j)}$  calculé après la reconstruction de la  $j^{eme}$  variable.

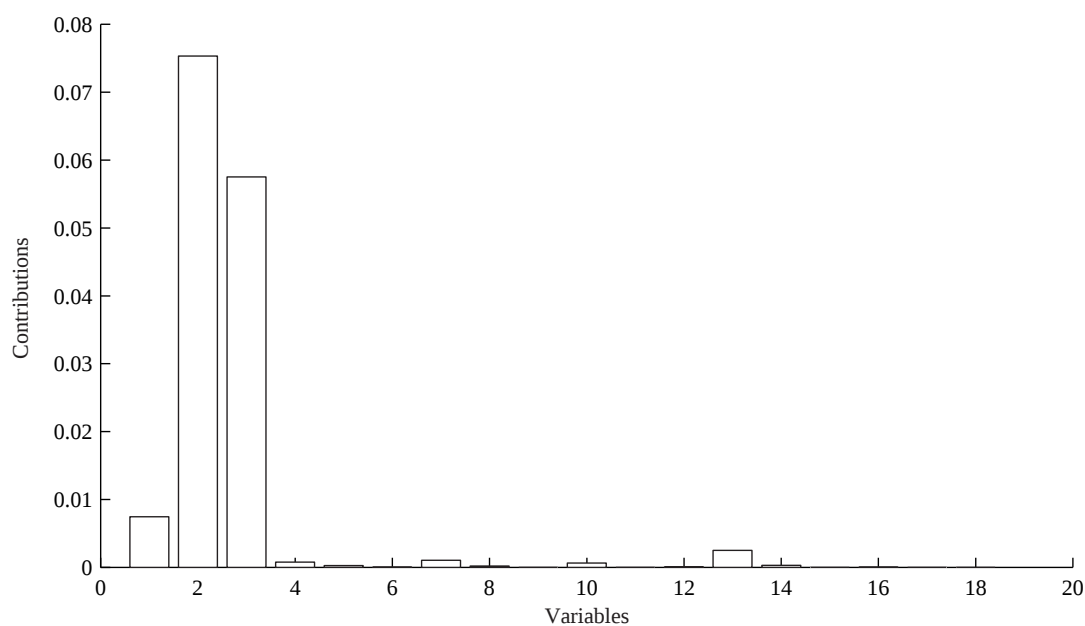


Figure 5.28 – Localisation de la variable  $v_2$  par calcul des contributions des variables à l'indice  $D_1$ .

L'application de la procédure de détection et de localisation proposée a permis de détecter et de localiser, à chaque fois, la variable incriminée en utilisant l'indice  $D_i$ . La reconstruction de la variable localisée a permis d'obtenir une estimation très satisfaisante de cette dernière et

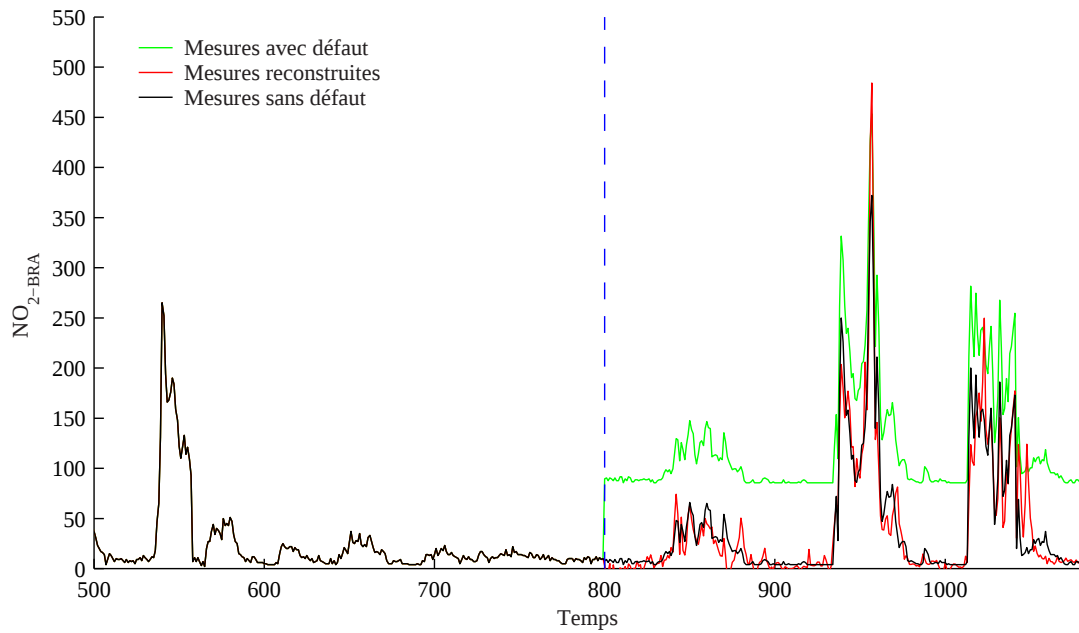


Figure 5.29 – Reconstruction de la variable en défaut du dioxyde d’azote de Brabois  $\text{NO}_2\text{-BRA}$ .

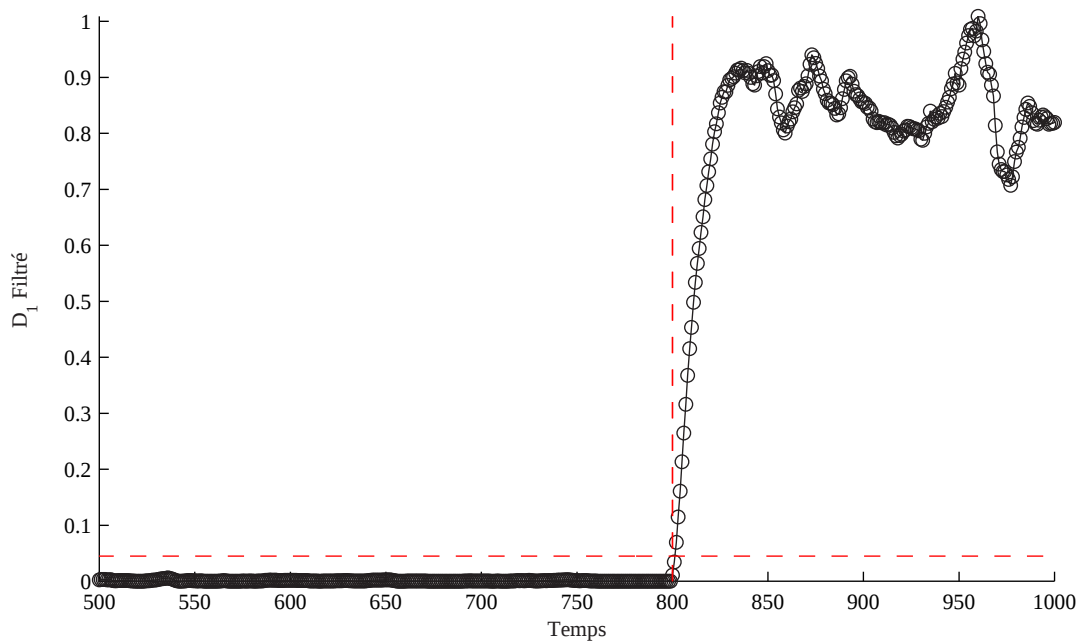
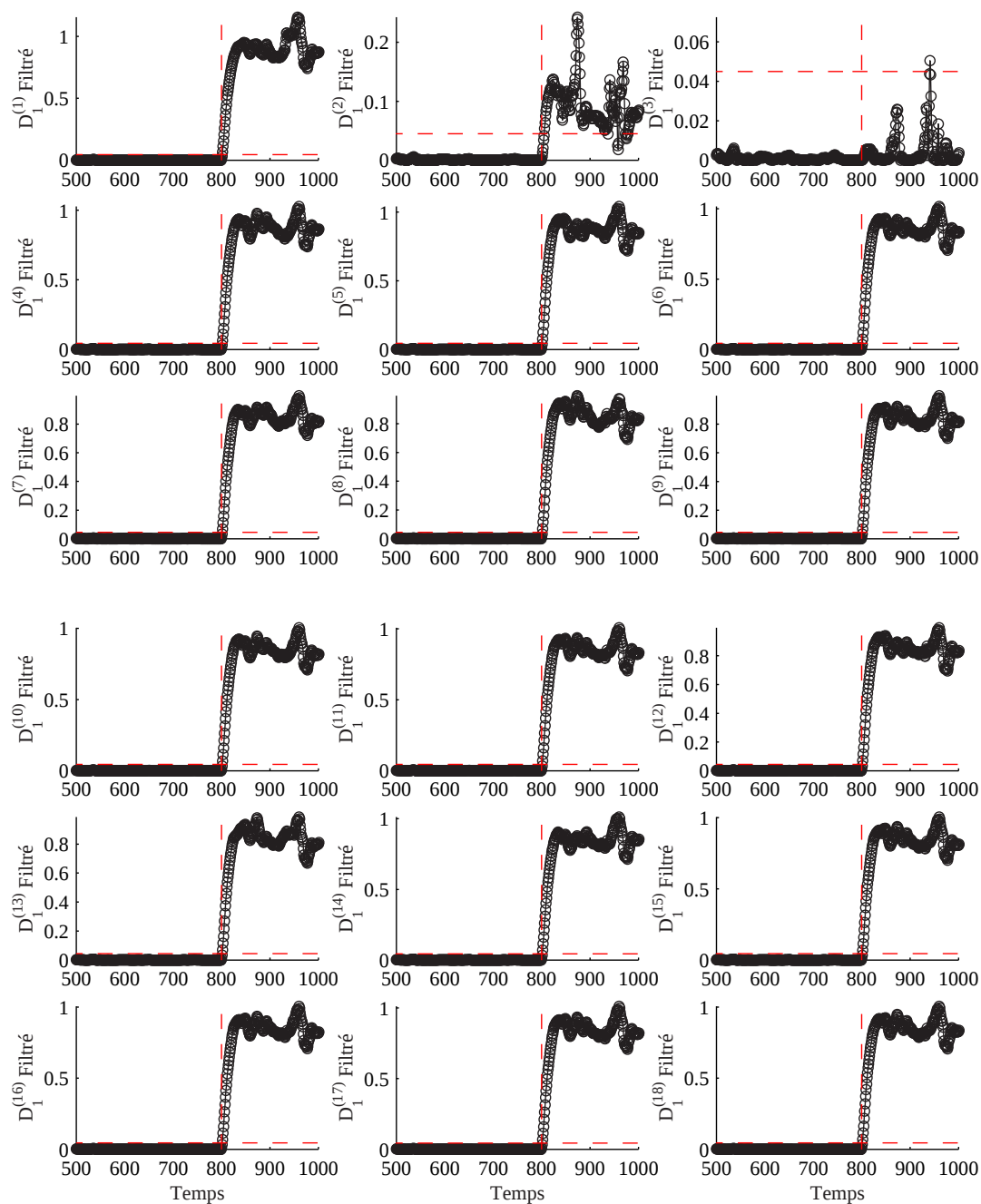


Figure 5.30 – Evolution de l’indice  $D_1$  filtré en présence d’un défaut affectant la variable  $v_3$ .

d’éliminer l’effet du défaut.

De plus, on arrive même à détecter et à localiser des défauts affectant certains oxydes d’azote et à reconstruire également la variable en défaut comme on l’a constaté avec des défauts simulés sur les oxydes d’azote mesurés sur le site de Brabois.

Figure 5.31 – Localisation de la variable  $v_3$  par reconstruction des 18 variables.

Ainsi, nous avons montré la faisabilité pour une partie du réseau et pour une période de temps de quelques jours. Les résultats obtenus nous encouragent à approfondir l'étude pour finaliser une procédure de diagnostic.

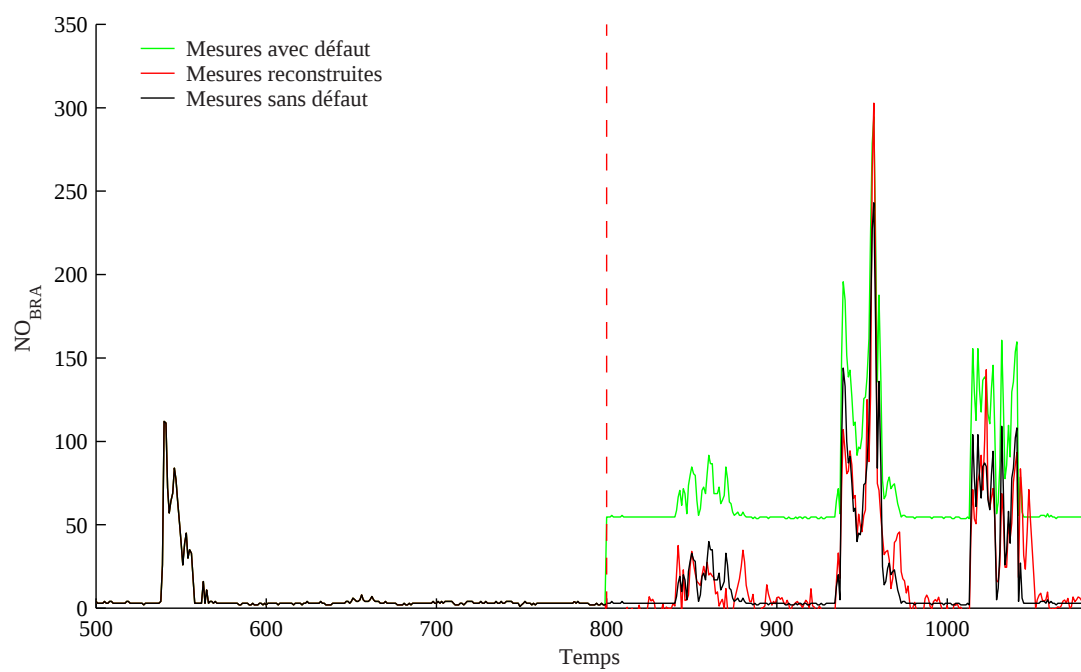


Figure 5.32 – Reconstruction de la variable en défaut  $NO_{BRA}$  de Brabois.



# 6

## Conclusions et Perspectives

### 6.1 Conclusion

Dans le domaine du diagnostic, des méthodes basées sur le concept de redondance de l'information ont été développées. Leur principe repose généralement sur un test de cohérence entre un comportement observé du processus fourni par les capteurs et un comportement prévu fourni par une représentation mathématique du processus. Ces méthodes de redondance analytique nécessitent un modèle du système à surveiller. Ce modèle comprend un certain nombre de paramètres dont les valeurs sont supposées connues lors du fonctionnement normal. Le prix à payer est toutefois l'élaboration d'un modèle mathématique aussi complet que possible dont la qualité est primordiale pour l'obtention d'un système de détection performant. Cela est tout à fait envisageable lorsqu'il s'agit des systèmes de petite dimension (nombre de variable à surveiller). En revanche, pour les systèmes de taille plus importante, établir de telles relations mathématiques entre les variables paraît moins immédiat. De plus, le choix de modèle pour la génération de résidus est arbitraire et il se peut très bien que les corrélations entre certaines variables ne soient pas prises en compte. Comme alternative, les méthodes basées sur l'analyse en composantes principales (ACP), sont très intéressantes pour la mise en évidence des corrélations linéaires significatives entre les variables du processus sans formuler de façon explicite le modèle du système. Un de nos objectifs a donc été d'utiliser l'analyse en composantes principales pour la détection et la localisation de défauts de capteurs.

Avant tout, nous avons présenté le principe de l'analyse en composantes principales (ACP). Cette méthode permet d'étudier les relations qui existent entre les variables et d'identifier les dépendances entre les observations multivariées afin d'obtenir une représentation compacte de ces dernières. Dans ce travail l'ACP est utilisée comme une technique de modélisation des relations entre les différentes variables du processus. L'estimation des paramètres du modèle ACP est effectuée par estimation des valeurs et vecteurs propres de la matrice des corrélations des données. Cependant, pour la détermination de la structure du modèle, il faut déterminer le nombre de composantes à retenir dans ce modèle (nombre de vecteurs propres). Pour cette raison, plusieurs critères de sélection du nombre de composantes ont été présentés. Notre choix s'est porté sur le critère utilisant le principe de reconstruction en minimisant la variance de l'erreur d'estimation des variables reconstruites, ce critère est très intéressant pour le diagnostic

car il tient compte de la redondance entre les mesures. L'approche de reconstruction permet également une sélection des variables à surveiller. Ainsi, une procédure permettant à la fois de déterminer les variables à surveiller et le nombre de composantes à retenir dans le modèle ACP peut être appliquée sur les données collectées sur le système ou le processus à surveiller.

Une fois le modèle identifié (détermination à la fois des variables à surveiller et du nombre de composantes à retenir ainsi que l'estimation des vecteurs propres de la matrice de corrélation des données), la procédure de détection et localisation de défauts peut être effectuée par génération des indicateurs de défauts (résidus) en comparant le comportement observé du processus donné par les variables mesurées et le comportement prévu donné par le modèle ACP. La plupart des méthodes de détection à base de l'ACP utilisent la statistique *SPE* (erreur quadratique) et la statistique  $T^2$  de Hotelling pour la détection de défauts. Cependant, le *SPE* est un test global qui cumule les erreurs de modélisation présentes sur chaque résidu et la statistique  $T^2$  n'est pas efficace car les conditions d'utilisation de cette statistique sont rarement vérifiées d'autant plus que cette statistique est calculée à partir des premières composantes principales qui ne représentent en aucun cas des résidus et donc ne représente pas un indice de détection. Deux exemples de simulation ont été présentés pour illustrer l'approche proposée et pour montrer les problèmes rencontrés avec les approches classiques.

Pour ces raisons, nous avons proposé un nouvel indice de détection basé sur les dernières composantes principales. Cet indice est défini par la somme successive des carrés des dernières composantes. Ainsi, l'expression générale de ce test est similaire à celle d'un *SPE* défini dans des sous-espaces inférieurs à celui du *SPE*. De même, les seuils de détection peuvent être calculés avec un raisonnement similaire à celui du calcul du seuil pour le *SPE* en utilisant les dernières valeurs propres.

Cet indice présente l'avantage d'être moins sensible aux erreurs de modélisation et plus sensible aux défauts que le *SPE* et de détecter des défauts avec des amplitudes inférieures à celles que peut détecter un *SPE*.

Concernant le problème de localisation, plusieurs approches de localisation basées sur l'analyse en composantes principales ont été présentées. Nous avons présenté dans un premier temps la méthode basée sur la structuration des résidus ou les résidus primaires sont transformés en des résidus secondaires avec une certaine propriété de localisation. Les résidus structurés obtenus par cette méthode dans le cas de l'ACP présentent de nombreuses fausses alarmes de plus que les conditions d'utilisation de cette approche avec les résidus utilisés pour le calcul de l'indice proposé ne sont pas vérifiées.

La deuxième approche que nous avons présentée ici, utilise des bancs de modèles comme dans le cas des observateurs. Dans cette catégorie d'approches on trouve trois méthodes : la méthode utilisant des modèles ACP partielles (ACP avec un nombre réduit de variables), la méthode utilisant le principe de reconstruction et la méthode utilisant le principe d'élimination.

Dans le cadre de la méthode utilisant des ACP partielles, nous avons montré avec un exemple simple qu'avec cette méthode on ne sait pas s'il est possible d'élaborer un modèle ACP réduit a priori, d'autant plus qu'elle utilise l'indice *SPE* comme indicateur de défaut. La deuxième méthode présentée est basée sur le principe de reconstruction qui consiste à suspecter qu'un capteur est défaillant et à reconstruire la valeur de sa mesure en se basant sur le modèle ACP déjà calculé et les mesures des autres capteurs. Comme cette approche utilise initialement l'indice *SPE* nous avons proposé de l'adapter à notre indice de détection. La procédure de localisation proposée permet de localiser les défauts simples et peut être utilisée pour la localisation de défauts multiples. La dernière méthode présentée dans la catégorie des méthodes utilisant des bancs de modèles est l'approche par élimination. Cette méthode est très proche de la méthode par reconstruction ou un ensemble d'indicateur de défaut sont générés en éliminant à chaque fois une

variable de l'ensemble des variables à surveiller bien sûr en éliminant, de la matrice du modèle ACP, la colonne correspondant à la variable éliminée. Comme toutes les autres approches, cette dernière utilise l'indice *SPE*. Pour cette raison nous avons adapté cette approche pour l'appliquer à l'indice de détection proposé.

La dernière méthode de localisation basée sur l'analyse en composantes principales est la méthode utilisant le calcul des contributions des variables aux indices de détection. A partir de l'expression de l'indice que nous proposons, nous avons proposé deux définitions des contributions des variables à l'indice proposé. La première définition est basée sur le fait que cet indice est un *SPE* particulier et de ce fait on peut exploiter la définition des contributions dans le cas du *SPE*. La deuxième définition que nous proposons vient du fait que l'indice proposé peut être calculé également à partir des dernières composantes principales. Ainsi, un calcul des contributions aux dernières composantes ayant subies une variation significative peut être utilisé pour cette définition. Dans les deux définitions, la variable ayant la plus grande contribution est considérée comme la variable incriminée.

Dans le troisième chapitre de cette thèse, nous avons présenté l'extension de l'ACP dans le cas non-linéaire (ACP<sub>NL</sub>). Nous avons présenté les principales méthodes pour le calcul du modèle ACP<sub>NL</sub>. La méthode la plus connue est la méthode des courbes principales, cependant cette méthode donne une estimation des variables originelles et des composantes principales non-linéaires mais ne permet pas d'avoir un modèle de représentation pour une estimation en ligne. Les méthodes les plus utilisées, dans le domaine du diagnostic, font appel aux réseaux de neurones. La méthode la plus utilisée et la plus connue utilise un réseau à cinq couches. L'apprentissage d'un tel réseau est très coûteux en terme de temps de calcul en plus du problème de convergence. Pour cette raison, nous avons proposé une nouvelle méthode pour le calcul du modèle ACP<sub>NL</sub> combinant la méthode des courbes principales et les réseaux à fonction de base radiale (RBF). Avec cette approche l'apprentissage se ramène à un problème de régression d'où un gain considérable en temps de calcul. La deuxième proposition dans ce chapitre concerne la détermination du nombre de composantes à retenir dans le modèle ACP<sub>NL</sub>. En exploitant le critère utilisant le principe de reconstruction dans le cas linéaire, nous avons proposé une extension de ce critère dans le cas non-linéaire. Un exemple de simulation a été introduit pour illustrer les méthodes proposées.

Dans le dernier chapitre nous avons présenté l'application des méthodes proposées dans le cas linéaire à la détection et la localisation de défauts de capteurs d'un réseau de surveillance de la qualité de l'air en Lorraine. L'étude menée est une étude de faisabilité et donc nous nous sommes limités à des périodes de mesures de quelques jours. Avec l'indice proposé on arrive à détecter et à localiser des défauts sur les variables d'ozone et même sur certains oxydes d'azote qui sont des variables plus localisées que l'ozone et donc plus difficile à modéliser. Les résultats obtenus sont très encourageants à approfondir l'étude pour finaliser une procédure de diagnostic.

## 6.2 Perspectives

Dans le deuxième chapitre nous avons proposé d'étendre l'approche de localisation, utilisant le principe de reconstruction avec l'indice de détection, aux défauts multiples. Toutefois, la procédure proposée n'est pas automatique et nécessite l'intervention d'un opérateur pour déterminer les variables à reconstruire. Une perspective directe de ce travail consiste à automatiser cette procédure par analyse de la matrice des derniers vecteurs propres.

Dans le cas de l'ACP non linéaire, l'approche proposée utilise un algorithme des courbes principales choisi sans tenir compte des autres algorithmes des courbes principales. Ainsi, une étude des différents algorithmes d'extraction des courbes principales est indispensable pour la

sélection d'une approche qui sera combinée avec les réseaux RBF. Concernant le problème de recherche de structure, dans le cas de l'ACP<sub>NL</sub>, deux problèmes peuvent être traités. Le premier concerne la détermination du nombre de neurones dans la couche cachée et la détermination du nombre de neurones de la couche de sortie qui représente le nombre de composantes principales non-linéaires. Pour la détermination du nombre de neurones de la couche cachée plusieurs méthodes peuvent être utilisés comme l'algorithme OLS (orthogonal Least square) par exemple. Cependant, pour le nombre de composantes à retenir dans le modèle ACP<sub>NL</sub>, l'algorithme que nous avons proposé peut être utilisé. Ainsi, une procédure d'optimisation plus complète peut être proposée. De plus concernant la procédure de sélection du nombre de composantes non-linéaire il reste à déterminer les condition de reconstruction et de convergence de cet algorithme.

En ce qui concerne l'application traitée dans le dernier chapitre, quelques propositions sont faites pour améliorer la qualité du modèle ACP. Une proposition à moyen terme consiste à utiliser différents modèles ACP pour différentes zones de fonctionnement, vu que les différents polluants présentent un comportement cyclique jour-nuit, ils peuvent être modélisés par une ACP utilisant des lots de données où chaque lot correspond à une période de fonctionnement. Ainsi, une utilisation de la MWPCA (Multi-Way PCA) peut être exploitée dans ce cas. Une autre solution consiste à exploiter le formalisme multi-modèles pour la modélisation en utilisant des multi-ACP comme une première extension vers un modèle non-linéaire.

L'application de l'approche utilisant les courbes principales et les réseaux RBF pour la modélisation des différents polluants, peut être envisagée une fois la procédure d'identification de structure du réseau achevée.

## Références bibliographiques



## Bibliographie

- [1] Rapport de l'académie des sciences. (1993). Ozone et propriétés oxydantes de la troposphère, Technique et documentation. Lavoisier.
- [2] Akaike H. (1974). Information theory and an extension of the maximum likelihood principle. In Proceedings 2nd International Symposium on Information theory, Petrov and Caski Eds., pp. 267-281.
- [3] Besse P. et Ferré L. (1993). Sur l'usage de la validation croisée en analyse en composantes principales. *Revue de Statistique Appliquée*, XLI (1), pp. 71-76.
- [4] Boudaoud N., Cherfi Z. (2000). Maîtrise statistique des processus multivariés : Avantages et limites des différentes approches sur les cartes de contrôle multivariées. *Journal Européen des Systèmes Automatisés*, vol. 34, pp. 379-390.
- [5] Box G. E. P. (1954). Some theorems on quadratic forms applied in the study of analysis of variance problems : Effect of inequality of variance in one-way classification. *The Annals of Mathematical Statistics*, vol. 25, pp. 290-302.
- [6] Brunet M., Jaume D., Labarrère M., Rault A., Vergé M. (1990). Détection et diagnostic de pannes - Approche par modélisation. *Traité des Nouvelles Technologies, série Diagnostic et Maintenance*, Hermès, Paris.
- [7] Conlin A. K., Martin E. B. and Morris A. J. (2000). Confidence limits for contribution plots. *Journal of Chemometrics*, vol. 14, pp. 725-736.
- [8] Cybenko G. (1989). Approximation by superposition of sigmoidal function. *Mathematics of Control Signal and Systems*, vol. 2, pp. 303-314.
- [9] Delicado P. (1998). Principal curves and principal oriented point. Technical Report, N° 309, Departament d'Economia Empresa, Univesitat Pompeu Fabra.
- [10] Der R., Steinmetz U., Balzuweit G. (1998). Nonlinear principal component analysis. Technical Report at the Institut für Informatik, Universität Leipzig.
- [11] Diamantaras K. I., Kung S. Y. (1996). Principal component neural networks. Theory and applications. John Wiley & Sons.
- [12] Dong D. and McAvoy T. J. (1994). Nonlinear principal component analysis - based on principal curves and neural networks. *Proceeding of the American Control Conference, ACC*.
- [13] Doymaz F., James C., Romagnoli J. A., Palazoglu A. (2001). A robust strategy for real-time process monitoring. *Journal of Process Control*, vol. 11, pp. 343-359.
- [14] Dunia R., Qin S. J. and Edgar T. F. (1996). Identification of faulty sensors using principal component analysis. *AIChE Journal*, vol. 42, N°. 10, pp. 2797-2812.
- [15] Dunia R., Qin S. J. and Edgar T. F. (1996). Multivariable process monitoring using nonlinear approaches. *Proceeding ACC, Seattle, Washington*, pp. 756-760.
- [16] Dunia R., Qin S. J. (1998). Joint diagnosis of process and sensor faults using principal component analysis. *Control Engineering Practice*, vol. 6, pp. 457-469.

- [17] Dunia R. and Qin S. J. (1998). A unified geometric approach to process and sensor fault identification and reconstruction : the unidimensional fault case. *Computers Chem. Engng.*, vol. 22, N°. 7-8, pp. 927-943.
- [18] Dunia R. and Qin S. J. (1998). A subspace approach to multidimensional identification and reconstruction. *AIChE Journal*, vol. 44, pp. 1813-1831.
- [19] Ferré L. (1995). Selection of components in principal component analysis : A comparison of methods. *Computational Statistics and Data Analysis*, pp. 669-682.
- [20] Frank P. M. (1990). Fault diagnosis in dynamic systems using analytical and knowledge - based redundancy - a survey and some new results. *Automatica*, vol. 26, N° 3, pp. 459-474.
- [21] Funahashi K. (1989). On the approximate realization of continuous mappings by neural networks. *Neural Networks*, vol. 2, pp. 183-192.
- [22] Gertler J. (1988). Survey of model-based failure detection and isolation in complex plants. *IEEE Control Systems Magazine*, vol. 8, N° 6, pp. 3-11.
- [23] Gertler J. (1991). Analytical redundancy methods in fault detection and isolation - survey and synthesis. *Proceeding of the IFAC Symposium on Fault Detection Supervision and Safety for Technical Process*, Baden Baden, Germany, pp. 9-22.
- [24] Gertler J., and Monajemy R. (1993). Generating directional residuals with dynamic parity equations. *IFAC Symposium on Fault Detection Supervision and Safety for Technical Process*, Sydney, Australia, pp. 507-512.
- [25] Gertler J. and McAvoy T. (1997). Principal component analysis and parity relations - A strong duality. *IFAC Conference SAFEPROCESS*, Hull, UK, pp. 837-842.
- [26] Gertler J., Weihua L., Yunbing H. and McAvoy T. (1998). Isolation enhanced principal component analysis. *3<sup>rd</sup> IFAC Workshop on On-line Fault Détection and Supervision in the Chemical Process Industries*, Lyon, June 4-5, France.
- [27] Gertler J. (1998). *Fault detection and diagnosis in engineering systems*. Marcel Dekker Editions, New York.
- [28] Gnanadesikan R. (1977). *Methods for statistical data analysis of multivariate observations*. Wiley Editions, New York.
- [29] Hansen C. (1992). *Regularization tools - A MATLAB package for analysis and solution of discrete ill-posed problems*. Technical University of Denmark. Disponible sur le site [http ://www.imm.dtu.dk/ pch/](http://www.imm.dtu.dk/pch/).
- [30] Gasso K., Harkat M. F., Mourot G. et Ragot J. (1998). Elaboration de modèles pour la validation, la crédibilité des données et la prédiction du taux d'ozone et de NO<sub>2</sub> en région lorraine. *Rapport N°1*, 52 pages.
- [31] Harkat M.F., Mourot G., Ragot J. (2000). Sensor failure detection of air quality monitoring network. *IFAC Symposium on Fault Detection, Supervision and Safety for Technical Processes, SAFEPROCESS'2000*, Budapest, Hungary, June 14-16.
- [32] Harkat M.F., Mourot G., Ragot J. (2000). Détection de défauts de capteurs d'un réseau de surveillance de la qualité de l'air. *Conférence Internationale Francophone d'Automatique, CIFA'2000*, Lille, France.
- [33] Harkat M.F., Mourot G., Ragot J. (2001). Sensor failure detection and Isolation of air quality monitoring network. *4th International Conference on Acoustical and Vibratory Surveillance Methods and Diagnostic Techniques*, Compiègne, France.

- [34] Harkat M.F., Mourot G., Ragot J. (2001). Détection et localisation de défauts de capteurs d'un réseau de surveillance de la qualité de l'air. 2ème Colloque Automatique et Environnement, A&E'2001, St Etienne, France.
- [35] Harkat M. F., Mourot G., Ragot J. (2002). Différentes méthodes de localisation de défauts basées sur les dernières composantes principales. Conférence Internationale Francophone d'Automatique, CIFA'02, Nantes, France.
- [36] Harkat M. F., Mourot G., Ragot J. (2003). Variable reconstruction using RBF-NLPCA for process monitoring. IFAC Symposium on Fault Detection, Supervision and Safety for Technical Processes, SAFEPROCESS'2003.
- [37] Harkat M. F., Mourot G., Ragot J. (2003). Nonlinear PCA combining principal curves and RBF-Networks for process monitoring. Soumis au CDC'2003.
- [38] Hastie T. (1984). Principal curves and surfaces. PhD thesis, Stanford University.
- [39] Hastie T. and Stuetzle W. (1989). Principal curves. Journal of the American Statistical Association, vol. 84, N° 406, pp. 502-516.
- [40] Himes D. M., Storer R. H. and Georgakis C. (1994). Determination of the number of principal component for disturbance detection and isolation. Proceedings of ACC, Baltimore.
- [41] Hotelling, H. (1933). Analysis of a complex of statistical variables into principal components. Journal of Educational Psychology, vol. 24, pp. 417-441.
- [42] Hsieh W. W. (2001). Nonlinear principal component analysis by neural networks. Journal of Climate.
- [43] Huang Y. and Gertler J. (1999). Fault isolation by partial PCA and partial NLPCA. IFAC'99, 14th Triennial world congress, Beijing, P. R. China, pp. 545-550.
- [44] Huang B. (2001). Process identification based on last principal component analysis. Journal of Process Control, vol. 11, pp. 19-33.
- [45] Isermann R. (1984). Process fault detection based on modeling and estimation methods - a survey. Automatica, vol. 20, pp. 387-404.
- [46] Jackson J. E., Mudholkar G. S. (1979). Control procedures for residuals associated with principal component analysis. Technometrics, vol. 21, N° 3.
- [47] Jia F., Martin E. B. and Morris A. J. (1999). Multiple sensor disturbance identification through principal component analysis. IFAC, 14<sup>th</sup> Triennial World Congress, Beijing, China.
- [48] Jolliffe I. T. (1986). Principal component analysis. Springer-Verlag, New York.
- [49] Juinghui C. and Jialin L. (1999). Mixture principal component analysis models for process monitoring. Industrial & Engineering Chemistry Research, vol. 39, pp. 1478-1488.
- [50] Kano M., Shinji H. and Hashimoto I. (2000). Contribution plots for fault identification based on the dissimilarity of process data. AIChE Annual Meeting, paper 255e, Los Angeles, CA, Nov. 12-17.
- [51] Kano M., Ohmo H., Shinji H. and Hashimoto I. (2001). New multivariate statistical process monitoring method using principal component analysis. Computers and Chemical Engineering, (in press), 2001.
- [52] Kégl B., Krzyzak A., Linder T. and Zeger K. (1999). Learning and design of principal curves. IEEE Transactions on Pattern analysis and Machine Intelligence.
- [53] Kerling M. (1999). Optimizing the multilayer perceptron - Problems, tools and strategies. Proc. of Eufit'99, Aachen, Germany.

- [54] Kerschen G. and Golinval J. C. (2002). Nonlinear generalisation of principal component analysis : from a global to a local approach. *Journal of Sound and Vibration*.
- [55] Kourti T., and MacGregor J. F. (1996). Recent Developments in Multivariate SPC Methods for Monitoring and Diagnosing Process and Product Performance. *Journal Of Quality Technology*, vol. 28, N° 4, pp. 409-428.
- [56] Kramer M. A. (1991). Nonlinear principal component analysis using autoassociative neural networks. *AIChE Journal*, vol. 37, N° 2, pp. 233-243.
- [57] Kresta J. V., MacGregor J. F. and Marlin T. E. (1991). Multivariate statistical monitoring of process operating performance. *Can. J. Chem. Eng.*, vol. 69, N° 1, pp. 35-47.
- [58] Kreyszig E. (1959). *Differential geometry*. University of Toronto Press, Toronto.
- [59] Kung S. Y., Diamantaras K. I. (1990). A neural network learning algorithm for adaptive principal component extraction (APEX). *Proceeding of the IEEE International Conference on Acoustics Speech and Signal Processing*, pp. 861-864.
- [60] LeBlanc M. and Tibshirani R. (1994). Adaptive principal surfaces. *Journal of American Statistical Association*, vol. 89, N° 425, pp. 53-64.
- [61] Levenberg K. (1944). A method for the solution of certain nonlinear problem in least squares. *Quart. Appl. Math.*, vol. 2, pp. 164-168.
- [62] MacGregor J. F. and Kourti T. (1995). Statistical process control of multivariate process. *Control Engineering Practice*, vol. 3, N° 3, pp. 403-414.
- [63] MacGregor J. F., Kourti T. and Nomikos P. (1996). Analysis, monitoring and fault diagnosis of industrial processes using multivariate statistical projection methods. *IFAC, 13th Triennial Word Congress*, pp. 145-150, San Francisco, USA.
- [64] Malthouse E. C., Mah R. S. H., and Tanhane A. C. (1995). Some theoretical results on nonlinear principal component analysis. *Proceeding of the ACC, Seattle, Washington*, pp. 744-748.
- [65] Maquin D., Ragot J. (2000). *Diagnostic des systèmes linéaires*. Hermès, Collection pédagogique d'automatique, Paris.
- [66] Marquardt D. W. (1963). An algorithm for least-squares estimation of nonlinear parameters. *J. Soc. Appl. Math.*, vol. 11, pp. 431-441.
- [67] Monahan A. H. (2000). *Nonlinear principal component analysis of climate data*. PhD thesis, British Columbia University.
- [68] Miller P., Swanson R. E., and Heckler C. E. (1998). Contribution plots : A missing link in multivariate quality control. *Applied Mathematics and Computer Science*, vol. 8, N° 4, pp. 775-792.
- [69] Morris A. J. and Martin E. B. (1997). Process performance monitoring and fault detection through multivariate statistical process control. *IFAC Conference SAFEPROCESS'97*, Hull UK, pp. 1-14.
- [70] Nandakishore A. and Leen T. K. (1997). Dimension reduction by local PCA. *Neural Computation*, vol. 9, pp. 1493-1510.
- [71] Nomikos P. and MacGregor J. (1995). Multivariate SPC Charts for Monitoring Batch Processes. *Technometrics*, vol. 37, N° 1, pp. 41-59.
- [72] Oja E. (1982). A simplified neuron model as a principal component analyzer. *Journal Math. Biology*, vol. 15, pp. 267-273.

- [73] Oja E. (1989). Neural networks, principal components, and subspaces. *International Journal of Neural Systems*, vol. 1, pp. 61-68.
- [74] Orr M. J. (1996). Introduction to radial basis function networks. Rapport technique disponible sur le site <http://www.anc.ed.ac.uk/mjo/rbf.html>
- [75] Pearson, K. (1901). On lines and planes of closest fit to systems of points in space. *The London, Edinburgh and Dublin Philosophical Magazine and Journal of Science*, vol. 6, pp. 559-572.
- [76] Qin S. J., Hongyu Y. and Dunia R. (1997). Self validating inferntial sensors with application to air emission monitoring. *Industrial & Engineering Chemistry Research*, vol. 36, pp. 1675-1685.
- [77] Qin S. J. and Dunia R. (1998). Determining the number of principal components for best reconstruction. *Proc. of the 5-th IFAC Symposium on Dynamics and Control of process Systems*, pp. 359-364, Corfu, Greece.
- [78] Qin S. J., Weihua L., Yue H. (1999). Recursive PCA for adaptive process monitoring. *IFAC'99, 14th Triennial World Congress*, Beijing, China.
- [79] Qin, S. J. and Weihua L. (1999). Detection, identification and reconstruction of faulty sensors with maximized sensitivity. *AIChE Journal*, vol. 45, N° 9, pp. 1963-1976.
- [80] Raich A. and Cinar A. (1996). Process disturbance diagnosis by statistical distance and angle measures, In *Proceedings of IFAC Congress*, San Francisco, vol. N, pp. 283-288.
- [81] Rissanen J. (1978). Modelling by shortest data description. *Automatica*, vol. 14, pp. 465-471.
- [82] Roweis S. T. (1998). EM algorithms for PCA and SPCA. In Michael I. Jordan, Michael J. Kearns, and Sara A. Sola, Editors, *Advances in Neural Information Processing Systems*, vol. 10, MIT Press.
- [83] Rubner J., Tavan P. (1989). A self-organizing network for principal component analysis. *Europsychics Lettre*, vol. 20, pp. 693-698.
- [84] Sanger T. D. (1989). Optimal unsupervised learning in a single-layer linear feedforward neural network. *Neural Networks*, vol. 2, pp. 459-473.
- [85] Shewhart, W. A. (1931). *Economic control of quality of manufactured product*. D. Van Nostrand, New York, NY.
- [86] Stewart G. W. (1973). *Introduction to matrix computations*. Academic Press, London.
- [87] Simoglou A., Martin E. B., Morris, A. J. (1997). Multivariate statistical process control in chemicals manufacturing. *IFAC Conference SAFEPROCESS'97*, Hull UK.
- [88] Spearman, C. (1904). General intelligence objectively determined and measured. *American Journal of Psychology*, vol. 15, pp. 201-293.
- [89] Stork C. L., Veltkamp D. J. and Kowalski B. R. (1997). Identification of multiple sensor disturbances during process monitoring. *Analytical Chemistry*, vol. 69, N° 24, pp. 5031-5036.
- [90] Staroswiecki M., Cassar J. P. (1996). Approche structurelle pour la conception des systèmes de surveillance. *Ecole d'Été d'Automatique*, Grenoble, 2-6 septembre.
- [91] Spearman, C. (1904). General intelligence objectively determined and measured. *American Journal of Psychology*, vol. 15, pp. 201-293.
- [92] Tan S. and Mavrovouniotis M. L. (1995). Reduction data dimentionaliti through optimizing neural network inputs. *AIChE Journal*, vol. 41, N° 6, pp. 1471-1480.

- [93] Teppola P., Mujunen S., Minkkinen P., Puijola T., Pursiheimo P. (1998). Principal component analysis, contribution plots and feature weights in the monitoring of sequential process data from a paper machine's wet end. *Chemometrics and Intelligent Laboratory Systems*, vol. 44, pp. 307-317.
- [94] Thissen U., Willem J. M., Lutgarde M. C. B. (2001). Nonlinear process monitoring using bottle-neck neural networks. *Analytica Chimica Acta*, vol. 446, pp. 371-383.
- [95] Thurstone, L. L. (1947). *Multiple factor analysis*. Chicago : University of Chicago Press.
- [96] Tibshirani R. (1992). Principal curves revisited. *Statistics and computation*, vol. 2, pp. 183-190.
- [97] Tong H. and Crowe C. M. (1995). Detection of Gross errors in data reconciliation by Principal Component Analysis. *Process Systems Engineering*, vol. 41, N° 7.
- [98] Tong H. and Bluck D. (1998). An industrial application of principal component test to fault detection and identification. *3<sup>rd</sup> IFAC Workshop on On-line Fault Detection and Supervision in the Chemical Process Industries*, Lyon, France.
- [99] Valle S., Weihua L., and Qin S. J. (1999). Selection of the number of principal components : The variance of the reconstruction error criterion with a comparison to other methods. *Industrial & Engineering Chemistry Research*, vol. 38, pp. 4389-4401.
- [100] Van Huffel S. (1997). Recent advances in total least squares techniques and errors-in-variables modeling. *Proceedings of the Second International Workshop on Total Least Squares and Errors-in-Variables Modeling*. Edited by Sabine Van Huffel, Katholieke Universiteit Leuven, Leuven-Heverlee, Belgium.
- [101] Verbeek J. J., Vlassis N. and Kröse B. (2000). A k-segments algorithm for finding principal curves. *IAS Technical Report Series*, nr. IAS-UVA-00-11, University of Amsterdam.
- [102] Voss B. (1999). A simulation study on nonlinear principal component analysis. *Technical Report 42/99 SFB 475*, Department of Statistics, University of Dortmund.
- [103] Wax M., Kailath T. (1985). Detection of signals by information criteria. *IEEE Trans. Acoust. Speech Signal Process.* ASSP-33, pp. 387-392.
- [104] Webb A. R. (1996). An approach to nonlinear principal component analysis using radially symmetric kernel functions. *Statist. Comput.*, vol. 6, pp. 159-168.
- [105] Webb A. R. (1999). A loss function to model selection in nonlinear principal components. *Neural Networks*, vol. 12, pp. 339-345.
- [106] Weihua L., Sirish S. (2002). Structured residual vector-based approach to sensor fault detection and isolation. *Journal of Process Control*, vol. 12, pp. 429-443.
- [107] Wenfu K., Storer R. H. and Georgakis C. (1995). Disturbance detection and isolation by dynamic principal component analysis. *Chemometrics and intelligent laboratory systems*, vol. 30, pp. 179-196.
- [108] Westerhuis J. A., Gurden S. P., Smilde A. K. (2000). Generalized contribution plots in multivariate statistical process monitoring. *Chemometrics and Intelligent Laboratory Systems*, vol. 51, pp. 95-114.
- [109] Westerhuis J. A., Gurden S. P., and Smilde A. K. (2000). Standardized Q-statistic for improved sensitivity in the monitoring of residuals in MSPC. *Journal of Chemometrics*, vol. 14, pp. 335-349.
- [110] Willsky A. S. (1976). A survey of design methods for failure detection in dynamic systems. *Automatica*, vol. 12, pp. 601-611.

- 
- [111] Wilson D. J. H., Irwin G. W. (1999). RBF principal manifolds for process monitoring. *IEEE Transaction on Neural Networks*, vol. 10, N° 6, pp. 1424-1434.
  - [112] Wise B. M. (1991). *Adapting Multivariate Analysis for Monitoring and Modeling of Dynamic Systems*. Ph.D. Dissertation, University of Washington, Seattle.
  - [113] Wise B. M. and Gallagher N. B. (1996). The process chemometrics approach to process monitoring and fault detection. *Journal of Process Control*, vol. 6, N° 6, pp. 329-348.
  - [114] Wold S., Esbensen K., Geladi P. (1987). Principal component analysis. *Chemom. Intell. Lab. Syst.*, vol. 2, pp. 37-52.
  - [115] Yue H. H. and Qin S. J. (2001). Reconstruction-based fault identification using a combined index. *Industrial & Engineering Chemistry Research*, vol. 40, pp. 4403-4414.





## Résumé

Les travaux présentés dans ce mémoire sont axés sur la détection et la localisation de défauts en utilisant l'analyse en composantes principales (ACP).

Dans le premier chapitre les principes fondamentaux de l'analyse en composantes principales linéaire sont présentés. L'ACP est utilisée pour la modélisation des processus en fonctionnement normal.

Dans le deuxième chapitre le problème de détection et localisation de défauts par ACP linéaire est abordé. A partir de l'analyse des indices de détection classiques, un nouvel indice de détection de défaut basé sur les dernières composantes principales a été développé. Pour la localisation de défauts, les méthodes classiques, utilisant par exemple le principe de reconstruction ou encore le calcul des contributions à l'indice de détection, ont été adaptées à l'indice de détection proposé.

Le troisième chapitre est consacré à l'ACP non-linéaire (ACP<sub>NL</sub>). Une extension de l'ACP pour des systèmes non-linéaires, combinant l'algorithme des courbes principales et les réseaux RBF, est proposée. Pour la détermination du nombre de composantes à retenir dans le modèle ACP<sub>NL</sub>, une extension du critère basé sur la variance de l'erreur de reconstruction a été proposée.

Une application réalisée dans le cadre d'une collaboration avec le réseau de surveillance de la qualité de l'air en Lorraine AIRLOR fait l'objet du quatrième chapitre. Cette application concerne la détection et la localisation de défauts de capteurs de ce réseau en utilisant la procédure de détection et de localisation développée dans le cas linéaire.

**Mots-clé :** Diagnostic, Analyse en composantes principales, Détection et localisation de défauts de capteurs, Reconstruction de variables, Contribution des variables, Réseaux RBF, Qualité de l'air.

---

## Abstract

The aim of this thesis is to study the fault detection and isolation using principal components analysis (PCA).

In the first chapter the fundamental principles of linear principal component analysis are presented. PCA is used to model normal process behaviour.

In the second chapter the problem of fault detection and isolation based on linear PCA is tackled. On the basis of the analysis of the classical detection indices, a new fault detection index based on the last principal components is developed. For fault isolation, the classical methods, using for instance the reconstruction principle or the contribution calculation, are adapted for the proposed fault detection index.

The third chapter is focused on the nonlinear PCA. An extension of the PCA for nonlinear systems, combining principal curves algorithm and RBF networks, is proposed. For the determination of the number of principal components to be kept in the NLPCA model, we propose an extension of the unreconstructed variance criteria in the non-linear case.

Finally, an application, carried out in collaboration with air quality monitoring network in Lorraine AIRLOR, is presented in the fourth chapter. This application concerns the sensor fault detection and isolation of this network by using the fault detection and isolation procedure developed in the linear case.

**Keywords :** Diagnosis, Principal component analysis, Sensor fault detection and isolation, Variable reconstruction, Variable contribution, RBF Neural Networks, Air quality.