

X

TA169.6

GG8

UNIVERSIDAD DE LOS ANDES
FACULTAD DE INGENIERÍA
CONSEJO DE ESTUDIOS DE POSTGRADO
DIVISIÓN DE POSTGRADO
POSTGRADO EN AUTOMATIZACIÓN E INSTRUMENTACIÓN

DETECCIÓN Y DIAGNÓSTICO DE FALLAS
UTILIZANDO TÉCNICAS ESTADÍSTICAS
MULTIVARIABLES A PARTIR DE DATA
HISTÓRICA

WWW.BDIGITAL.ULA.VE

Por:

José Luis Gouveia Ramírez

Tutores:

Dr. Oscar Camacho

M.Sc. Delfina Padilla

Trabajo de Grado presentado como requisito parcial
para obtener el Grado de
MAGÍSTER SCIENTIAE EN AUTOMATIZACIÓN
E INSTRUMENTACIÓN

MÉRIDA, OCTUBRE de 2005

DONACION

SERBIULA
Tulio Febres Cordero

Licencia Creative Commons:
Atribución - No Comercial - Compartir Igual 3.0 Venezuela
(CC BY-NC-SA 3.0 VE)

AGRADECIMIENTOS

WWW.BDIGITAL.ULA.VE

A mis padres por darme todo el apoyo necesario para poder obtener esta meta.

A Andreina por motivarme en la aventura de realizar un postgrado y animarme en aquellos momentos donde la carga parecía muy pesada, love este logro también es tuyo.

A los profesores Oscar Camacho y Delfina Padilla por toda su colaboración y consejos.

A todos los profesores y personal del postgrado en Automatización e Instrumentación por estar siempre a la orden cuando los necesite.

A mis compañeros del postgrado y del LABIDAI, especialmente a Arnaldo, Ender y David por haber sido de gran apoyo en todas esas pequeñas pero importantes cosas del día a día.

Licencia Creative Commons:
Atribución - No Comercial - Compartir Igual 3.0 Venezuela
(CC BY-NC-SA 3.0 VE)

RESUMEN

En este trabajo se aplican algunas técnicas estadísticas multivariantes para realizar la detección, identificación y diagnóstico de fallas en procesos industriales. A partir de dos procesos simulados, el proceso de un tanque de reacción no isotérmico agitado continuamente (CSTR) y el Proceso Tennessee Eastman (TEP), se obtiene la data histórica en presencia de fallas típicas, de donde se genera un patrón que luego permite identificar una falla similar a este. El Análisis de Componentes Principales (PCA), el Análisis Discriminante de Fisher (FDA) y el Análisis Discriminante Generalizado (GDA) son algunas de las técnicas estadísticas que se emplean en este trabajo. Se demostró la aplicabilidad del GDA para clasificar fallas donde la data requiere de un clasificador no lineal. Por último, se propone un nuevo procedimiento completo para el estudio de las fallas que incluye, para la etapa de detección de fallas, el estadístico Hotelling y el método DISSIM en paralelo. Luego en la etapa de extracción de características, se propone el uso del Análisis Wavelet. La identificación de fallas se realiza por medio del Análisis de pares FDA, que permite seleccionar las variables más relacionadas con la falla. En la etapa del diagnóstico de fallas se propone el FDA donde la data requiere de un clasificador lineal y el GDA donde la data requiere de un clasificador no lineal.

Palabras claves: Detección, identificación y diagnóstico de fallas; Análisis de Componentes Principales (PCA); Análisis Discriminante de Fisher (FDA); Análisis Discriminante Generalizado (GDA); Análisis Wavelet; estadístico Hotelling; DISSIM.

ÍNDICE GENERAL

Agradecimientos.....	i
Resumen.....	ii
Índice General.....	iii
Lista de Figuras.....	vi
Lista de Tablas.....	viii

CAPÍTULO I

INTRODUCCIÓN.....	1
-------------------	---

CAPÍTULO II

MARCO TEORICO.....	12
2.1. Conceptos Generales.....	12
2.2. Objetivos del Sistema de Detección y Diagnóstico de Fallas.....	14
2.3. Características de un SDDF.....	14
2.4. Diagnóstico de fallas por medio de clasificación de patrones.....	16
2.5. Métodos para la Detección y Diagnóstico de Fallas.....	18
2.5.1. Método para la detección de fallas basado en el Análisis de Componentes Principales (PCA).....	18
2.5.2. Método para la detección de fallas basado en la disimilaridad de la data del proceso (DISSIM).....	20
2.5.3. Análisis Discriminante de Fisher (FDA).....	22
2.5.4. Análisis Wavelet.....	25
2.5.5. Análisis Discriminante Generalizado (GDA).....	27
2.5.6. Análisis de pares FDA.....	30
2.5.7. Análisis Discriminante.....	32
2.6. Procedimiento propuesto para el estudio de las fallas.....	33

CAPÍTULO III

CASOS DE ESTUDIO.....	35
3.1 Descripción del proceso de un tanque de reacción no isotérmico agitado continuamente (CSTR).....	35
3.1.1 Sistema de Control.....	39
3.1.2 Descripción de las Fallas.....	41
3.1.3 Variables medidas del proceso CSTR.....	43
3.1.4 Generación de la data histórica.....	44
3.2 Descripción del Tennessee Eastman Process (TEP).....	47
3.2.1 Sistema de Control.....	48
3.2.2 Descripción de las Fallas.....	49
3.2.3 Variables del TEP.....	49
3.2.4 Generación de la data histórica.....	51

CAPÍTULO IV

RESULTADOS Y DISCUSIÓN.....	54
4.1 Detección de Fallas del CSTR utilizando el estadístico Hotelling (T^2).....	54
4.2 Identificación de fallas del CSTR utilizando las direcciones de falla en pares FDA.....	56
4.3 Clasificación de las Fallas del CSTR utilizando Análisis Discriminante (AD).....	57
4.4 Clasificación de las Fallas del CSTR utilizando FDA y AD.....	59
4.5 Clasificación de las Fallas del CSTR utilizando Wavelet, FDA y AD.....	61
4.6 Clasificación de las Fallas del CSTR utilizando Wavelet, GDA y AD.....	62
4.7 Detección de Fallas del TEP utilizando el estadístico Hotelling (T^2)...	64

4.8 Detección de Fallas del TEP utilizando el método DISSIM.....	65
4.9 Identificación de fallas del TEP utilizando las direcciones de falla en pares FDA.....	67
4.10 Clasificación de las Fallas del TEP utilizando Análisis Discriminante (AD).....	69
4.11 Clasificación de las Fallas del TEP utilizando FDA y AD.....	70
4.12 Clasificación de las Fallas del TEP utilizando Wavelet, FDA y AD.....	71
4.13 Clasificación de las Fallas del TEP utilizando Wavelet, GDA y AD.....	72
CAPÍTULO V	
CONCLUSIONES.....	74
REFERENCIAS BIBLIOGRÁFICAS.....	76

LISTA DE FIGURAS

1.1. Métodos de Diagnóstico.....	3
2.1. Esquema de un SDDF y fuentes de posibles fallas.....	13
2.2. Fallas separadas por planos en R^m	16
2.3. Esquema típico de un sistema de clasificación de patrones.....	17
2.4. Objetivo del FDA.....	24
2.5. Izquierda: Data de las Clases <i>A</i> y <i>B</i> en el espacio original 2-D.	
Derecha: Data proyectada en un espacio característico 3-D.....	27
2.6. Artificio Kernel. El producto punto en el espacio característico	
kernel <i>F</i> es definido implícitamente.....	28
2.7. Izquierda: Data Iris de Fisher aplicando FDA.	
Derecha: Data Iris de Fisher aplicando GDA.....	30
2.8. Algoritmo propuesto para la detección y diagnóstico de fallas.....	34
3.1. Proceso del CSTR con control cascada.....	36
3.2. Comparación de la variable “Temperatura del reactor (T)” para la condición	
normal de operación y la Falla 4p (Bias en la medición de temperatura T)..	46
3.3. Comparación de la variable “Flujo del refrigerante (Q_c)” para la condición	
normal de operación y la Falla 4p (Bias en la medición de temperatura T)..	46
3.4. Diagrama de proceso del TEP.....	47
3.5. Comparación de la variable “Componente A (corriente 6) (% molar)” para	
la condición normal de operación y la Falla 1.....	52
3.6. Comparación de la variable “Flujo Alimentación de A (corriente 1)” para la	
condición normal de operación y la Falla 1.....	53
4.1 Gráfica de control T^2 para las fallas 4 y 21.....	55
4.2 Gráfica de contribución FDA para las fallas 10 y 12, data de entrenamiento.	56
4.3. Selección del número de autovectores <i>N</i> por medio de validación cruzada.	59
4.4 Gráfica de control de T^2 para la falla 2,3, 8 y 9.....	65

4.5 Gráfica de control del DISSIM para las fallas 2 y 9.....	66
4.6 Gráfica de contribución FDA para las fallas 4 y 6, data de validación y (td-td+100).....	68
4.7 Selección del número de autovectores N por medio de validación cruzada..	70

WWW.BDIGITAL.ULA.VE

LISTA DE TABLAS

3.1. Parámetros y condiciones de operación nominal.....	38
3.2. Parámetros de modelado de los transmisores del sistema de control.....	39
3.3. Presiones nominales y coeficientes de las válvulas.....	40
3.4. Parámetros de los controladores PI.....	41
3.5. Fallas y perturbaciones del CSTR.....	42
3.6. Variables medidas del Proceso CSTR.....	43
3.7. Parámetros de modelado de los transmisores.....	44
3.8. Desviación estándar del ruido.....	45
3.9. Fallas del TEP.....	49
3.10. Variables medidas del TEP.....	50
3.11. Variables manipuladas del TEP.....	51
4.1. Tiempos de retardo y confiabilidad del T^2 para la data de entrenamiento y validación.....	55
4.2 Variables Identificadas para la data de entrenamiento.....	57
4.3. Error de clasificación utilizando AD para diferentes periodos de tiempo con la data de entrenamiento y validación.....	58
4.4. Error de clasificación utilizando FDA+AD para diferentes periodos de tiempo con la data de entrenamiento y validación.....	60
4.5. Error de clasificación utilizando Wavelet+FDA+AD para el periodo de tiempo correspondiente a (td-1801) con la data de validación.....	61
4.6. Error de clasificación utilizando Wavelet+FDA+AD para diferentes periodos de tiempo con la data de validación.....	62
4.7. Error de clasificación de subgrupos de fallas, utilizando SV, Wavelet, FDA, GDA y AD, para el periodo de tiempo correspondiente a (td-td+100) con la data de validación.....	63
4.8. Tiempos de retardo y confiabilidad del T^2 para la data de validación.....	64

4.9. Tiempos de retardo y confiabilidad del DISSIM para la data de validación.	66
4.10 Variables Identificadas para la data de validación.....	68
4.11. Error de clasificación utilizando AD con la data de validación.....	69
4.12. Error de clasificación utilizando FDA+AD con la data de validación.....	71
4.13. Error de clasificación utilizando Wavelet+FDA+AD para el periodo de tiempo correspondiente a (td-800) con la data de validación.....	72
4.14. Error de clasificación del subgrupo de fallas, utilizando SV, FDA, GDA y AD, para el periodo de tiempo correspondiente a (td-td+100) con la data de validación.....	73

WWW.BDIGITAL.ULA.VE

Capítulo I. Introducción

Detección y Diagnóstico de Fallas es un área importante en ingeniería de procesos, debido a que interviene directamente en la rentabilidad y seguridad de los procesos de una planta. La pronta detección y diagnóstico de una falla, mientras el proceso se encuentra aun en una región controlable, permite evitar paradas de planta no programadas, reduciendo el tamaño de las pérdidas económicas y posibles situaciones de riesgo. Para tener una idea del tamaño de las pérdidas económicas que sufre la industria debido al efecto de fallas que no han sido detectadas y corregidas a tiempo, sólo en la industria petroquímica americana, un estimado de aproximadamente 20 mil millones de dólares se pierden anualmente y el costo es mucho mayor cuando incluimos otras industrias como la farmacéutica, química, generación de energía, etc., que en conjunto generan pérdidas alrededor de 27 mil millones de dólares anuales. Y por si fuera poco grandes accidentes en plantas químicas y petroquímicas han ocurrido en los últimos años en la Union Carbide's Bhopal en India, Petroleum's Piper Alpha y en la refinería Kuwait Petrochemical's Mina Al-Ahmedi, traducándose en grandes pérdidas económicas y humanas. [1]

Hoy en día se ha hecho popular en las empresas contar con un Sistema de Manejo de Situaciones Anormales (ASMS), el cual incluye tres etapas principales como son: la temprana detección de la condición de falla, el diagnóstico que determina la causa y por último la corrección de la falla para volver a colocar al proceso en su estado normal y seguro de operación. Nuestra principal empresa del país, PDVSA, en su locación de Occidente cuenta con su propio ASMS, conformado por un sistema experto que tiene el propósito de detectar condiciones anormales en el proceso de producción, específicamente en pozos, estaciones de flujo, sistema de

recolección/distribución de gas, plantas de compresión, etc. y alerta a los usuarios (ingeniero de optimización, ingeniero de gas, ingeniero de producción y operadores SCADA) acerca de las probables causas y recomendaciones para las acciones correctivas pertinentes.

Las necesidades, anteriormente mencionadas, hacen que el estudio de Detección y Diagnóstico de Fallas sea uno de los temas más interesantes y más estudiados en los últimos años y a medida que los procesos aumentan cada día en tamaño, integración y complejidad, es necesario investigar y desarrollar nuevas técnicas que brinden a los operadores de las plantas la mayor cantidad de información acerca de las situaciones de falla que se presentan en los procesos y de esta manera, ayudar en la toma de decisión de las acciones correctivas.

Algunas de las técnicas clásicas que fueron implementadas en los primeros desarrollos de sistemas de detección de fallas están basadas en la *redundancia física* y en la *verificación de valores límites*. La *redundancia física* requiere de la instalación de sensores redundantes, la detección se realiza comparando las salidas de los sensores, y por medio de un procedimiento de votación se diagnostica cual de los sensores esta fallando. También se puede realizar instalando sensores especiales que puedan medir alguna cantidad física indicadora de la presencia de una falla, tal como vibración, sonido, etc. La aplicabilidad de este método esta limitada al costo extra y el espacio adicional requerido, así como el hecho de que no puede detectar fallas que afecten a todos los instrumentos redundantes de la misma forma, como por ejemplo, efectos de temperatura, falla del suministro eléctrico, etc. Sin embargo, se utilizan comúnmente en sistema de control de aviones, vehículos espaciales y plantas nucleares donde la seguridad es crítica. En la *verificación de valores límites*, los datos provenientes del proceso son comparados con límites preestablecidos, al sobrepasarse el límite se indica la presencia de falla, generalmente hay dos límites el primero sirve de alarma y el segundo dispara mecanismos automáticos de emergencia. Los

inconvenientes que presenta este método es que para evitar que el sistema de detección de fallas no se dispare ante perturbaciones normales del proceso, los límites deben fijarse conservadoramente. En la mayoría de los casos no suministra suficiente información para identificar la causa de la falla y las consecuencias de la misma. [2]

El avance de los sistemas de computación ha permitido que nuevos métodos más eficientes puedan ser utilizadas para la detección y diagnóstico de fallas, estos se pueden clasificar en métodos basados en modelos cualitativos, modelos cuantitativos y basados en data histórica. Del trabajo de Venkatasubramanian [1] se obtiene una clasificación de estos métodos diagnóstico, mostrada en la Figura 1.1.

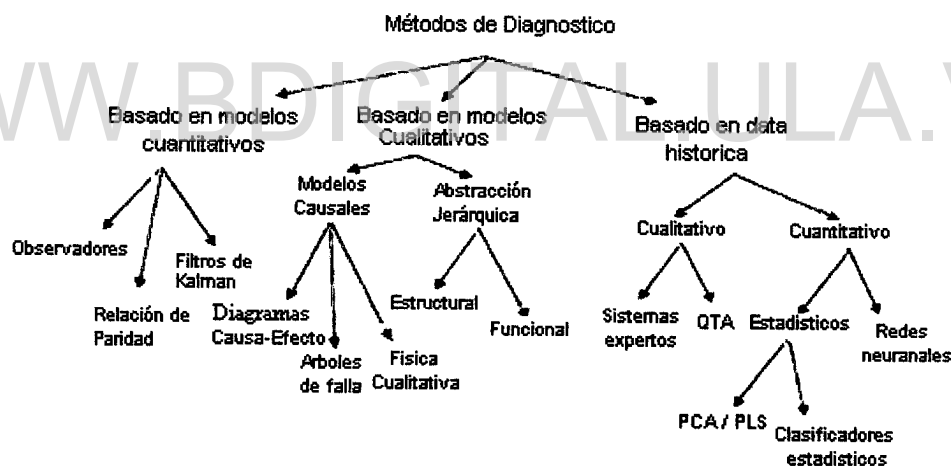


Figura. 1.1. Métodos de Diagnóstico

Los **métodos basados en modelos cuantitativos** se apoyan en el concepto de redundancia analítica, desarrollado por Willsky [3], donde se realiza una comparación de las salidas de los instrumentos que pueden presentar falla y valores esperados obtenidos por medio de un cálculo analítico de la variable, utilizando un modelo que

relaciona las variables del proceso. Este modelo puede ser obtenido matemáticamente a partir de las leyes físicas que rigen los componentes del proceso o modelos obtenidos experimentalmente de caja gris o negra, asumiendo linealidad en el proceso alrededor del punto de operación para facilitar la obtención de los modelos. De dicha comparación se generan residuos que, dependiendo si su valor es cercano a cero, no hay falla o si sobrepasa el umbral estamos en presencia de una falla.

Los métodos basados en modelos cuantitativos se clasifican en: observadores de diagnóstico, relaciones de paridad, filtros de kalman y estimación de parámetros. En los *observadores de diagnóstico* se desarrollan un grupo de observadores, cada uno de ellos es sensible a un subgrupo de fallas e insensible al resto de ellas. Si una falla ocurre todos los observadores que son insensibles a dicha falla presentan residuos muy pequeños, caso contrario al observador sensible a la falla que presenta residuos de gran magnitud, realizando de esta manera el diagnóstico. En el caso de *relaciones de paridad* la idea principal es chequear la paridad o consistencia de la salida del modelo de la planta con la salida de los sensores y con las entradas conocidas del proceso. Idealmente, en estado estacionario, el valor de las relaciones de paridad es cero. Aprovechando la capacidad de los *filtros de Kalman* como algoritmo recursivo para la estimación de estados puede ser realizado el diagnóstico de fallas monitoreando el error de predicción. Por último, en los métodos basados en *estimación de parámetros* se obtiene un modelo del proceso, mediante identificación cuyos parámetros son comparados con los del proceso sin falla. La diferencia entre dichos parámetros determina la presencia de una falla. Un ejemplo de esta técnica se encuentra en el trabajo de Garces [4], donde el diagnóstico se realiza a partir de identificación en línea con funciones de Laguerre, obteniéndose los coeficientes del modelo paramétrico de la posible falla, estos coeficientes se comparan utilizando el coeficiente de correlación de cada falla con patrones que fueron identificados previamente. Otro trabajo que aplica esta técnica, es el realizado por Sanchez [5], que implementa un sistema utilizando estructuras de transición, donde en una primera

etapa en línea, se realiza la detección de la falla si se sobrepasa un umbral del valor absoluto del error entre la salida del proceso y un modelo del proceso sin falla. El diagnóstico se realiza estableciendo el grado de pertenencia de la salida del proceso a cada una de los modelos de fallas, por medio de la prueba t-student. La obtención de dichos modelos se realiza previamente a partir de identificación de sistemas, conformando una librería con las fallas principales y secundarias del proceso del tipo ARX.

La desventaja que presentan estos métodos es que, en muchos casos, los procesos son muy complejos, multivariantes, no lineales y no es fácil obtener un modelo del proceso que sea lo suficientemente exacto que sirva de comparación para determinar la existencia de una falla, limitándolos a procesos sencillos.

En los **métodos basados en modelos cualitativos** la relación entre las entradas y salidas del sistema son expresadas en términos de funciones cualitativas. Estas funciones cualitativas se forman a partir de hechos y reglas de la descripción de la estructura y el comportamiento dinámico del proceso, el cual es caracterizado por síntomas o valores cualitativos tales como ON-OFF, valores límites, conjuntos difusos de valores tales como alto, pequeño, poco, etc. El primer método desarrollado fue el *sistema experto*, el cual consiste en un gran conjunto de reglas del tipo if-then-else que comparan una condición cualitativa del proceso con una base de conocimiento, para, de esta manera, obtener conclusiones que permiten detectar la falla, igual como lo haría una persona experta. Presenta el inconveniente que no tiene ningún entendimiento de las leyes físicas que rigen el sistema, por lo que falla al presentarse una nueva condición que no este presente en la base de conocimiento.

En los *diagramas de causa-efecto* existen nodos que representan eventos o variables y flechas que representan la relación entre los nodos, la dirección de las flechas van desde los nodos de causa a los nodos de efecto, correspondiendo cada nodo a la

variación de una variable del estado estacionario. Otro método que es usado para el análisis de las características cualitativas del proceso es el de *árboles de falla*, el cual realiza de forma deductiva la evaluación y análisis de las fallas, localizando y estableciendo su causa. La raíz del árbol de falla es un evento de falla del sistema que se denomina evento tope o cabecera. Las subfallas que conducen al evento cabecera forman los nodos del siguiente nivel del árbol, conectándose mediante puertas lógicas AND, OR, etc., expandiéndose en niveles posteriores siguiendo una estructura jerárquica. [1]

El problema de los métodos basados en modelos cualitativos es que no tienen ningún entendimiento de las leyes físicas del proceso, que trae como consecuencia que el método falle cuando aparezca una nueva condición no definida en la base de conocimiento.

En los **métodos basados en data histórica** la detección y diagnóstico de falla se realiza mediante el procesamiento de una gran volumen de data histórica, la cual se encuentra almacenada en las bases de datos de los sistemas de control SCADA y Sistemas de Control Distribuido utilizados por muchas industrias hoy en día. Hay varias maneras en las cuales la data puede ser transformada y presentada como un conocimiento *a priori* para el sistema de diagnóstico. Una es conocida como Extracción de Características (Feature Extraction). El proceso de extracción puede ser tanto cualitativo como cuantitativo [1]. Para el caso de la extracción cualitativa es común el uso de los sistemas expertos, lógica difusa y en el caso de la extracción cuantitativa tenemos las redes neuronales, el Análisis Wavelet, los métodos estadísticos como por ejemplo el Análisis de Componentes Principales (PCA) y el Análisis Discriminante de Fisher (FDA), etc.

Las redes neuronales del tipo feed forward han sido utilizadas satisfactoriamente en la detección y clasificación de fallas como se muestra en trabajo de Zhou [6]. Las Redes

Neuronales se comportan como un clasificador de patrones no lineal, debido a la posibilidad de seleccionar la función de transferencia de cada neurona y los pesos de conexión, contribuyendo al comportamiento no lineal de la red. Aportan, además, una importante solución a aquellos procesos donde es difícil obtener un modelo matemático preciso. En el trabajo de Rios [7] se presenta un método para la detección y diagnóstico de fallas en sistemas no lineales, basado en la reconstrucción en los modos de las fallas. Para reconstruir las fallas se utilizan modelos inversos que se obtienen por el diseño de redes neuronales, las cuales permiten aproximar el patrón de comportamiento de las salidas del sistema a las fallas.

El Análisis de Componentes Principales (PCA) es una técnica estadística multivariable que ha sido aplicada exitosamente en la detección y diagnóstico de fallas. En el trabajo de Padilla [8] el PCA es utilizado para el diagnóstico de fallas en motores de encendido por compresión y en el trabajo de Chiang [9] ha sido aplicada en la detección y diagnóstico de fallas de un proceso químico simulado. El Análisis Discriminante de Fisher (FDA), se presenta como una técnica estadística lineal muy efectiva en la clasificación de fallas, en el trabajo de Chiang [9] ha sido aplicada en el diagnóstico de fallas de un proceso químico simulado.

Un ejemplo de aplicación del Análisis Wavelet para la detección y diagnóstico de fallas fue realizado por Guillen [10] desarrollando un método que utiliza la transformada discreta Wavelet (TDW). El diagnóstico se realiza generando una serie de patrones para cada tipo de falla a partir de data histórica del proceso, conformando el vector característico utilizando la TDW sobre los patrones. El reconocimiento del tipo de falla se realiza por medio de técnica de reconocimientos de patrones tales como distancia euclídeana y k-nearest neighbor. Existen, además, trabajos que presentan híbridos de técnicas cualitativas y cuantitativas como el de Ruiz [11] que implementa un sistema compuesto por una red neuronal y un sistema de lógica difusa,

y el de Zhou [12] donde se realiza la detección y clasificación de fallas basado en Árboles de clasificación y FDA.

Los métodos basados en data histórica tienen una ventaja importante con respecto a los anteriores debido a que no requieren desarrollar un modelo matemático preciso del proceso, ni conjuntos de reglas de decisión, que en muchos procesos complejos es una tarea titánica. Además, el desarrollo de la instrumentación digital, las redes industriales y los sistemas de control SCADA y Control Distribuido, permiten almacenar un gran volumen de data de los procesos industriales hoy en día. Esta data histórica puede ser aprovechada para desarrollar sistemas de detección y diagnóstico de fallas, debido a que posee información directa de cómo se comportó el proceso real ante fallas y perturbaciones en el pasado, pudiendo, a partir de esta información, generarse un patrón que permita identificar una condición similar. Debido a estas razones, en este trabajo fue enfocado en los métodos basados en data histórica, específicamente en aquellos que utilizan técnicas estadísticas multivariantes para la realización de la detección, identificación y diagnóstico de fallas.

El Análisis de Componentes Principales (PCA) es una técnica estadística multivariante que ha sido aplicada exitosamente en la detección y diagnóstico de fallas [8,9,13]. PCA es usado en este trabajo para realizar la detección de fallas utilizando las gráficas de control Hotelling (T^2), que es una medida de variación entre un modelo PCA del proceso en estado normal de funcionamiento y en una condición de falla.

En los últimos años el Análisis Wavelet ha sido muy utilizado para analizar mediciones que contienen contribuciones a múltiples escalas y mediciones autocorrelacionadas, debido a su habilidad de comprimir características a múltiples escalas y aproximadamente decorrelacionar la data de muchos procesos estocásticos autocorrelacionados [14]. Los coeficientes Wavelet proveen información compacta

acerca de una señal en diferentes localizaciones de tiempo y frecuencia [15], hecho que permite también reducir la dimensionalidad de la data, pero en este caso en la dirección del tiempo a diferencia del FDA que lo hace en la dirección de las variables. Debido a su capacidad de reducir la dimensionalidad y decorrelacionar la data el Análisis Wavelet ha sido introducido en este trabajo en la etapa de extracción de características del diagnóstico de fallas.

Para el diagnóstico de fallas el uso del Análisis Discriminante de Fisher (FDA) se presenta como una técnica lineal muy efectiva en la clasificación de fallas, permite también realizar reducción de dimensionalidad [9], es decir reducir el número de variables que son obtenidas de la instrumentación del proceso en un grupo menor que son realmente las que poseen información importante de las fallas ocurridas, hecho que es muy común en los procesos industriales donde generalmente se miden y almacenan decenas de variables del proceso. Una limitante que posee el FDA es que es una técnica lineal, lo que quiere decir que las fallas son separadas usando planos, trayendo como consecuencia un alto error de clasificación para los casos en donde las fallas requieren de superficies no lineales o en el caso que las fallas se solapan unas con otras.

Para el caso de que la data requiere de un clasificador no lineal se propone el uso del Análisis Discriminante Generalizado (GDA), también llamado Análisis Discriminante Lineal Kernel (KLDA), el cual ha sido aplicado con éxito en la clasificación de patrones, como se muestra en el trabajo de Baudat [16] que comprueba su eficacia con la data de flores Iris de Fisher y en el trabajo de Koscor [17] donde muestra superioridad ante otras técnicas como el Análisis de Componentes Principales Kernel (KPCA) y el Análisis de Componentes Independientes Kernel (KICA) en el reconocimiento de voz. La comprobación de la eficacia de esta técnica en el diagnóstico de fallas es una de las motivaciones de este trabajo.

El objetivo principal de esta investigación es desarrollar un Sistema de Detección y Diagnóstico de Fallas que pueda ser implementado en todos aquellos procesos industriales en los que se tenga una base de datos históricos adecuada, utilizando para ello técnicas estadísticas multivariantes.

Para alcanzar este objetivo fue necesario cumplir con una serie de objetivos específicos, que incluían la selección de los procesos donde se evaluó el Sistema de Detección y Diagnóstico de Fallas propuesto. Un tanque de reacción no isotérmico agitado continuamente (CSTR) y el Proceso Tennessee Eastman (TEP), fueron seleccionados. De ambos procesos se estudiaron diferentes características tales como: fallas que presenta comúnmente, sistemas de control implementado e instrumentación presente. Luego se simuló el proceso CSTR en presencia de diferentes tipos de fallas y perturbaciones, utilizando como herramienta el programa Simulink de Matlab. Para el caso del TEP, la data se obtuvo directamente de [18]. De esta manera se generó la data histórica que permitió probar el desempeño del Sistema de Detección y Diagnóstico de Fallas propuesto, que incluye para realizar la detección, técnicas tales como el estadístico Hotelling T^2 y el método DISSIM. Luego en la etapa de extracción de características se propone el uso del Análisis Wavelet. El diagnóstico de fallas se realiza en dos etapas, un primer diagnóstico por medio de Análisis Discriminante de Fisher en conjunto con Análisis Discriminante que logra clasificar las fallas linealmente separables y una segunda etapa para aquellas fallas requieren de una segunda clasificación por medio del Análisis Discriminante Generalizado. La medición del desempeño del sistema de Detección y Diagnóstico de Fallas se realizó a partir del tiempo de retardo y la confiabilidad en la etapa de detección y del error de clasificación en la etapa del diagnóstico de las fallas.

El contenido de este trabajo está organizado de la siguiente manera. El **Capítulo I** presenta una introducción al tema de detección y diagnóstico de fallas, la justificación

del problema planteado, el estado del arte y los objetivos. El **Capítulo II** corresponde al marco teórico, compuesto por las generalidades relacionadas con la detección y diagnóstico de fallas, la descripción de los métodos de diagnóstico utilizados en este trabajo y el procedimiento propuesto. En el **Capítulo III** se analizan los casos de estudio que corresponde a dos procesos industriales el CSTR y el TEP seleccionados para probar el desempeño del procedimiento propuesto para la detección y diagnóstico de la fallas. El **Capítulo IV** muestra los resultados obtenidos y la discusión de los mismos. Por último en el **Capítulo V** se presentan las conclusiones de este trabajo.

WWW.BDIGITAL.ULA.VE

Capítulo II. Marco Teórico

En este capítulo se describen los conceptos y bases teóricas para poder realizar la detección y diagnóstico de fallas. En primer lugar son definidos los conceptos generales necesarios para abordar el tema, así como los objetivos y características del Sistema de Detección y Diagnóstico de Fallas (SDDF). Seguidamente se presentan las etapas para realizar el diagnóstico mediante clasificación de patrones a partir de data histórica. Luego se describen teóricamente los métodos utilizadas para la detección y diagnóstico de fallas. Finalmente se presenta el procedimiento completo propuesto para el estudio de las fallas, que incluye las etapas de detección, identificación y diagnóstico.

WWW.BDIGITAL.ULA.VE

2.1. Conceptos Generales

El termino *falla* es definido de manera general como aquel evento que es capaz de desviar una variable controlada o un parámetro asociado con el proceso de un rango deseado, este evento puede ser causado por desperfectos de los dispositivos de dicho proceso o *perturbaciones*. Para el caso de las *perturbaciones*, estas están asociadas con una variación de cantidades físicas o químicas que rigen el comportamiento del proceso y producidas por agentes externos (variación de la temperatura, caudal, presión, concentración, interferencia electromagnética, etc.) que fueron considerados constantes a la hora de realizar el diseño del sistema de control del proceso. A diferencia de las *fallas* que, a pesar de que pueden hacer variar una cantidad física o química, esta variación es producida por agentes internos del proceso, es decir ocurre por un desperfecto de un elemento del proceso o del sistema de control. Debido a que

las fallas y las perturbaciones tienen la misma consecuencia que es la desviación de la variable controlada, es tarea del SDDF detectar y diferenciar cuando ocurre cualquiera de ellas. En la Figura 2.1 se muestra esquemáticamente como un SDDF se acopla a un proceso que presenta un sistema de control. El SDDF obtiene información de las variables del proceso y de las señales de control que le permiten detectar y diagnosticar fallas, sin confundir estas con perturbaciones.

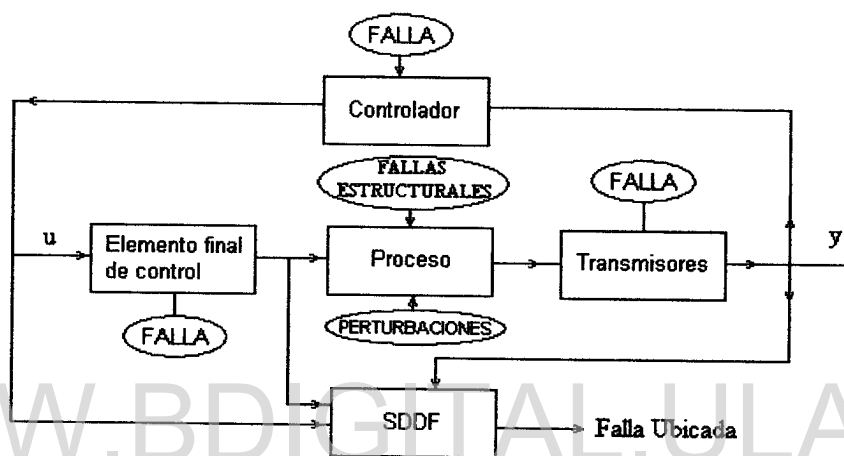


Figura 2.1. Esquema de un proceso con un SDDF y fuentes de posibles fallas.

En el esquema de la Figura 2.1 podemos observar los distintos tipos de fallas que pueden presentarse en un proceso y en el sistema de control. Las fallas en el proceso también llamadas estructurales, ocurren por malfuncionamiento de los equipos que conforman el proceso en si, por ejemplo: bombas, tuberías, calderas, mezcladores, quemadores, intercambiadores, motores, etc. En el sistema de control pueden presentarse una serie de fallas que afectan el desempeño del mismo, como son: fallas en transmisores o sensores que obtienen información de las variables del proceso (temperatura, flujo, presión, etc.), fallas en elementos finales de control tales como válvulas de control atascadas, que presentan juego, etc., y fallas propias del controlador. Por último el comportamiento del proceso también puede verse afectado por perturbaciones externas. Es deseable que todas estas posibles fallas y perturbaciones que puede presentarse en el proceso y en el sistema de control en un

momento dado sean detectadas, identificadas y diagnosticadas por el SDDF de manera rápida y confiable.

2.2. Objetivos del Sistema de Detección y Diagnóstico de Fallas

Como vimos anteriormente la detección, identificación y diagnóstico de fallas son los tres principales objetivos que un SDDF debe satisfacer al encontrarse el proceso en presencia de una falla o una perturbación. A continuación son definidos cada uno de estos objetivos:

- Detección de fallas: Establecimiento preciso del momento en que sucede la posible falla o perturbación en el sistema. La detección de fallas se realiza por medio de la generación de residuos, señales o síntomas que reflejen el suceso de una falla.

- Identificación de fallas: Determinación de las variables del proceso que se encuentran más relacionadas con la falla ocurrida. Lo que permite realizar más fácilmente el diagnóstico de fallas y ofrece información importante a los operadores de planta e ingenieros que a partir de su experiencia pueden localizar la causa de la falla.

- Diagnóstico de fallas: Clasificación exacta del tipo de falla o perturbación ocurrida. Se realiza seleccionando cual es la falla detectada de un conjunto de posibles fallas y perturbaciones que ocurren comúnmente en el proceso, usualmente por medio de técnicas de reconocimiento de patrones.

2.3. Características de un SDDF.

Además de la detección, identificación y diagnóstico de fallas, es deseable que el SDDF cumpla con una serie de características que aseguran su aplicabilidad en un

proceso real. Entre ellas tenemos la *rápida detección y diagnóstico*, que permite que los daños provocados por la falla al sistema sean menores, tomando medidas correctivas, preferiblemente, antes de que el proceso salga de su estado estable, evitando así situaciones de riesgo. Es importante que el SDDF cumpla con esta característica, sin dejar de ser *robusto* para no verse afectado por el ruido o incertidumbres y generar falsas alarmas o clasificaciones erradas.

Es conveniente que el SDDF pueda *reconocer nuevas fallas* y no clasificarlas equivocadamente como otra falla conocida o una operación normal del proceso. Esto no siempre es posible debido a que generalmente no se cuenta con una base de datos que contenga todas las posibles fallas o con un buen modelo de la dinámica del proceso que permita generar dicha data con facilidad, sin embargo siempre es importante tomar en cuenta la posibilidad de que el SDDF sea capaz de reconocer que una nueva falla ha ocurrido, así no se sepa a ciencia cierta de cual se trata. *Estimación del error de clasificación* es la capacidad de mostrar un valor del error de estimación o confiabilidad que se presenta en la detección y clasificación de fallas, siendo esta característica muy útil a la hora de tomar decisiones del tiempo y magnitud de las acciones correctivas a tomar por el personal de planta. Otra característica deseable pero no es sencilla de obtener, es la *adaptabilidad* a los posibles cambios de producción sin generar falsas alarmas.

Además de reconocer la causa es conveniente que el SDDF informe acerca del *efecto de la falla* a la estabilidad del proceso y de esta manera ayudar a la toma de decisiones por el personal de planta. Debido a lo complejo de los procesos actuales, y para un rápido y fácil desarrollo de los diagnósticos en tiempo real es conveniente que los *requerimientos de modelaje y computacionales* sean lo menos posible. Es difícil conseguir todas estas características en un mismo sistema por lo que se desea que posea la mayor cantidad.

2.4. Diagnóstico de fallas por medio de clasificación de patrones.

Actualmente gran cantidad de procesos industriales se encuentran automatizados, debido al desarrollo de la instrumentación, las redes industriales y los sistemas de control, esto permite almacenar un gran volumen de data de dichos procesos. En esta data histórica se encuentran almacenadas las mediciones de todas las variables del proceso en condiciones normales de operación y en presencia de fallas. A partir de esta información se puede generar un patrón, que permite, luego, determinar cuando una falla similar ocurra en el proceso. Para ello las diferentes fallas que se encuentran en la data histórica deben ser clasificadas en clases separadas. Geométricamente, las m variables ubican cada una de las fallas que contienen n observaciones en un espacio R^m , donde el número de ejes es igual al número de variables. Suponiendo que las fallas pueden ser separadas por planos como se muestra en la Figura 2.2, estos planos separadores definen los límites de la región de cada falla. Luego de que el sistema detecta una falla a partir de la data del proceso obtenida en línea, esta falla puede ser diagnosticada determinando en cual de las regiones de fallas esta se ubica, siempre y cuando la falla detectada se encuentre previamente en la base de datos históricos como se explico anteriormente.

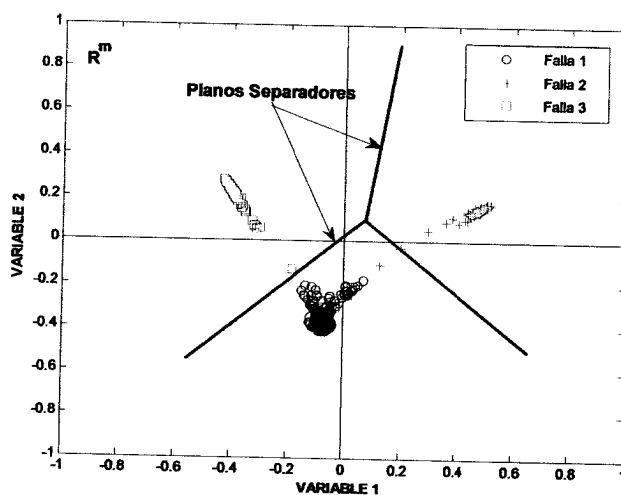


Figura. 2.2. Fallas separadas por planos en R^m

Un sistema típico de clasificación de patrones asigna un vector de observación a una clase específica por medio de los siguientes pasos: *extracción de características*, *cálculo discriminante* y *selección del máximo* como se muestra en la Figura 2.3. [9]. El objetivo de la etapa de *extracción de características* es incrementar la robustez del sistema de clasificación de patrones transformando las variables originales en un nuevo conjunto de variables que optimicen la separación entre las clases. Luego el *calculador discriminante* determina el valor de la **función discriminante** para cada clase, esta función cuantifica el grado de relación entre el vector de observación y cada una de las clases. La función discriminante indirectamente actúa como los planos separadores mostrados en la Figura 2.2. Dependiendo si esta función es lineal o no, el clasificador será también del tipo lineal o no lineal. Por último, se selecciona la clase con el máximo valor de la función discriminante en la etapa de *selección del máximo*.

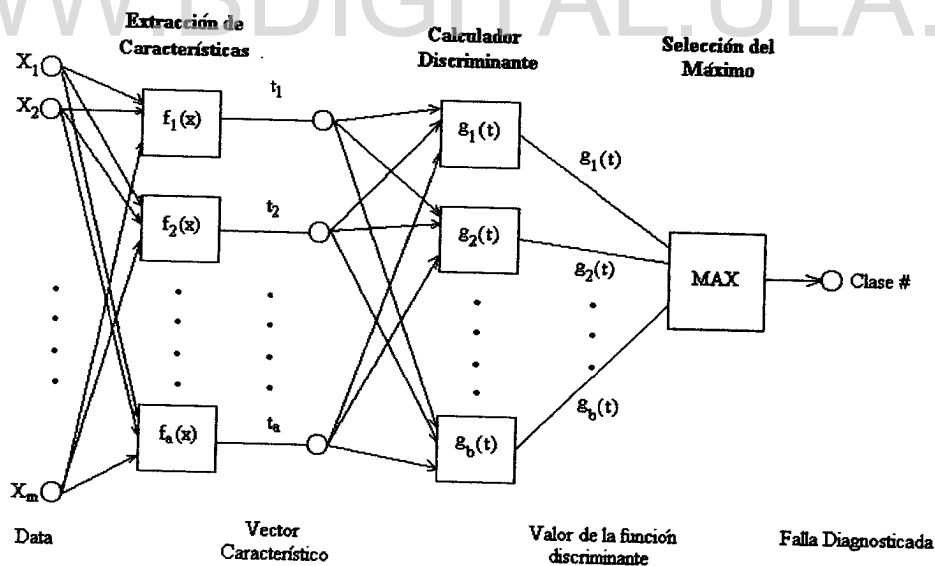


Figura 2.3. Esquema típico de un sistema de clasificación de patrones

Es importante destacar que las variables medidas del proceso deben permitir diferenciar los individuos en clases separables, por ello es necesario una correcta

selección de las variables a incluir en la base de datos, de hecho, generalmente las variables no siempre permiten separar las clases perfectamente, es común el hecho de que las fallas se solapen entre sí, existiendo un error de clasificación, o que se requieran hipersuperficies para poder realizar la separación de manera no lineal.

2.5. Métodos para la Detección y Diagnóstico de Fallas.

2.5.1. Método para la detección de fallas basado en el Análisis de Componentes Principales (PCA)

Una de las técnicas utilizadas en este trabajo para realizar la detección de fallas es el Análisis de Componentes Principales. PCA es una técnica lineal de reducción de dimensionalidad, óptima en términos de capturar la variabilidad de la data. Esta determina un conjunto de vectores ortogonales, que son combinaciones lineales de las variables originales, estos son ordenados de mayor a menor relevancia (contenido de información). Dada una matriz X de los datos originales, previamente centrada, de dimensión $n \times m$, donde m son las variables del proceso y n las observaciones, esta puede ser descompuesta en una suma de productos de vectores t_i (vector de coordenadas) y p_i (vector de pesos o ponderaciones) sumados a una matriz residual E , de la siguiente forma:

$$X = TP^T + E = \sum_{i=1}^a t_i p_i^T + E \quad (2.1)$$

La matriz P de los vectores p_i , forma una nueva base ortogonal para el espacio generado por X , donde los vectores p_i son los vectores propios de la matriz de covarianza de X definida como:

$$\text{cov}(X) = \frac{1}{n-1} X^T X \quad (2.2)$$

$$\text{cov}(X)p_i = \lambda_i p_i \quad (2.3)$$

Donde λ_i es el valor propio asociado con el vector propio p_i . Los modelos PCA generalmente son formados únicamente con un número a reducido de vectores propios que son los llamados componentes principales, debido a que estos son combinación lineal de las variables originales y juntos explican la mayoría de las variaciones de la data. Es, de esta manera, como se realiza la reducción de dimensionalidad de la data.

Cada uno de los vectores t_i es simplemente la proyección de X en el nuevo vector base p_i :

$$t_i = X p_i \quad (2.4)$$

El estadístico Hotelling (T^2) representa una medida de la variación entre el modelo PCA de la data de entrenamiento sin falla y una nueva data. Si el valor de T^2 supera un limite umbral estamos en la presencia de una falla, de esta manera podemos entonces realizar la detección de fallas en procesos multivariantes. El valor de T^2 para una nueva observación k es la suma de los valores normalizados al cuadrado y es definido como:

$$T^2(k) = t(k) \Lambda^{-1} t(k)^T \quad (2.5)$$

Donde Λ^{-1} es la inversa de la matriz diagonal de los valores propios asociados con los componentes principales retenidos para la data de entrenamiento sin falla. El límite umbral para T^2 se obtiene usando la distribución F .

$$T_{\alpha}^2 = \frac{a(n-1)(n+1)}{n(n-a)} F(a, n-a, \alpha) \quad (2.6)$$

Donde α es el nivel de significación, generalmente se elige un valor de 0,05 ó 0,01.

2.5.2. Método para la detección de fallas basado en la disimilaridad de la data del proceso (DISSIM)

Una de las técnicas utilizadas en este trabajo para realizar la detección de fallas es el DISSIM. Este método está basado en la idea que, un cambio de una condición de operación puede ser detectado monitoreando la distribución de la data del proceso debido a que esta refleja la correspondiente condición de operación [19]. Para evaluar las diferencias entre las distribuciones de los conjuntos de datos, es utilizado el método basado en la expansión de Karhunen_Loeve (KL).

Considerando dos matrices X_1 y X_2 . X_i consiste de n_i observaciones de m variables. Las matrices de covarianzas vienen dadas por:

$$R_i = \frac{1}{n_i} X_i^T X_i \quad (2.7)$$

La matriz de covarianza de la mezcla de ambos conjuntos de datos viene dado por R .

$$R = \frac{n_1}{n_1 + n_2} R_1 + \frac{n_2}{n_1 + n_2} R_2 \quad (2.8)$$

Basándose en el hecho de que la matriz de covarianza R puede ser diagonalizada por una matriz ortogonal P_o .

$$P_o^T R P_o = \Lambda \quad (2.9)$$

Las matrices originales X_i son transformadas en Y_i

$$Y_i = \sqrt{\frac{n_1}{n_1 + n_2}} X_i P_o \Lambda^{-1/2} \quad (2.10)$$

Las matrices de covarianza de las matrices transformadas

$$S_i = \frac{1}{n_i} Y_i^T Y_i = \frac{n_i}{n_1 + n_2} P^T R_i P \quad (2.11)$$

satisfacen las siguientes ecuaciones:

$$S_1 + S_2 = I \quad (2.12)$$

$$S_i w_j^{(i)} = \lambda_j^{(i)} w_j^{(i)} \quad (2.13)$$

Donde w_j y λ_j son los vectores propios y los valores propios, respectivamente.

Finalmente se define el índice D , mostrado a continuación, para evaluar la disimilaridad de los conjuntos de datos.

$$D = \frac{4}{m} \sum_{j=1}^m (\lambda_j - 0,5)^2 \quad (2.15)$$

Cuando los conjuntos de datos son similares entre si, el valor de D debe ser cercano a 0 y cercano a 1 cuando son muy diferentes.

Para aplicar DISSIM, un conjunto de datos de referencia y el límite umbral deben ser determinados previamente, a partir del siguiente procedimiento:

- 1) Adquirir un conjunto de datos cuando el proceso esta operando bajo condiciones normales de operación, normalizando cada variable de la matriz de datos.
- 2) Determinar el tamaño de la ventana de tiempo, w . Generando conjunto de datos con w observaciones de la data, moviendo la ventana de tiempo. Luego seleccionar un conjunto de datos de referencia.
- 3) Calcular el índice D y determinar el límite umbral mediante validación cruzada.

2.5.3. *Análisis Discriminante de Fisher (FDA).*

El Análisis Discriminante de Fisher (FDA) es una técnica lineal ampliamente usada para la clasificación de patrones, óptima en el término de maximizar la separación entre clases. El FDA fue utilizado en este trabajo en el diagnóstico de fallas, como una de las técnicas para realizar la etapa de extracción de características del sistema de clasificación de patrones. Una descripción matemática obtenida de [9] se muestra a continuación.

Dada una matriz X de dimensión $n \times m$, correspondiente a la data de entrenamiento para todas las fallas, donde la transpuesta de la i^{th} fila de X es el vector columna x_i , la matriz total de dispersión es:

$$S_t = \sum_{i=1}^n (x_i - \bar{x})(x_i - \bar{x})^T \quad (2.15)$$

donde $\bar{x}$ es el vector de media total, cuyos elementos corresponden a las medias de las columnas de X . Definiendo X_j como el conjunto de vectores x_i , los cuales pertenecen a la clase j , la matriz de dispersión-intra para la clase j es

$$S_j = \sum_{x_i \in X_j} (x_i - \bar{x}_j)(x_i - \bar{x}_j)^T \quad (2.16)$$

donde $\bar{x}_j$ es el vector de media para la clase j . Siendo c el número de clases o fallas, entonces la matriz de dispersión-intra-clases es

$$S_w = \sum_{j=1}^c S_j \quad (2.17)$$

y la matriz de dispersión-entre-clases es

$$S_b = \sum_{j=1}^c n_j (\bar{x}_j - \bar{x})(\bar{x}_j - \bar{x})^T \quad (2.18)$$

donde n_j es el número de observaciones en la clase j . La matriz total de dispersión resulta:

$$S_t = S_b + S_w \quad (2.19)$$

El objetivo del primer vector FDA φ_1 es maximizar la matriz de dispersión-entre-clases mientras se minimiza la matriz de dispersión-intra-clases (Figura.2.4):

$$\max \left\{ \frac{\text{dispersión - entre - clases}}{\text{dispersión - intra - clases}} \right\} = \max_{\varphi_1 \neq 0} \left\{ \frac{\varphi_1^T S_b \varphi_1}{\varphi_1^T S_w \varphi_1} \right\} \quad (2.20)$$

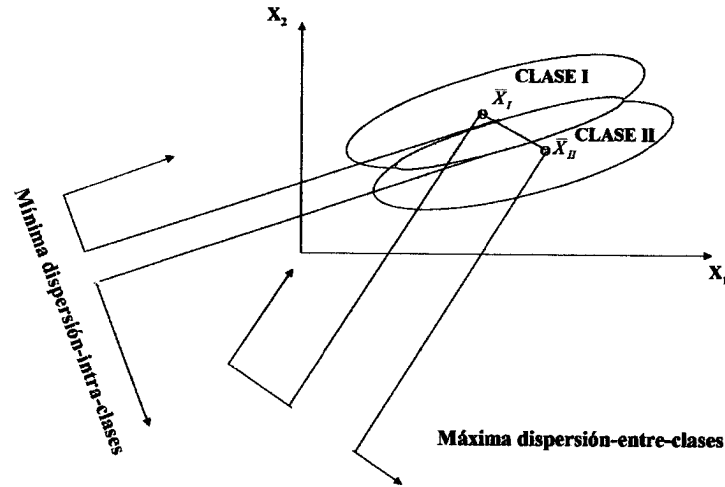


Figura 2.4. Objetivo del FDA

El segundo vector FDA se calcula para maximizar la dispersión-entre-clases mientras se minimiza la dispersión-intra-clases entre todos los ejes perpendiculares al primer vector FDA y de la misma manera para el resto de vectores FDA. Ha sido demostrado matemáticamente que los vectores FDA son iguales a los vectores propios ϕ_k del problema de valores propios generalizado:

$$S_b \phi_k = \lambda_k S_w \phi_k \quad (2.21)$$

Donde los valores propios λ_k indican el grado de separabilidad entre todas las clases proyectando la data en ϕ_k . El primer vector FDA es el vector propio asociado con el valor propio mayor, el segundo vector FDA es el vector propio asociado con el segundo valor propio mayor, así sucesivamente. La transformación lineal de la data original de m dimensiones a un espacio de N dimensiones es descrito por:

$$Z = \phi_N^T X \quad (2.22)$$

siendo N es el número de vectores propios.

La selección del número de vectores propios permite reducir la dimensionalidad de la data, seleccionando un número menor que el número de variables. En los casos que no se cuenta con suficiente data para estimar el vector de media y la matriz de covarianza de cada clase, al realizar una reducción de dimensionalidad se reduce el error de clasificación [9].

Comúnmente se determina el número apropiado de vectores propios mediante validación cruzada, donde la data se separa en dos conjuntos, una data de entrenamiento y una de validación, seleccionando el número de vectores propios que reduzca al mínimo el error de clasificación de la data de validación.

2.5.4. *Análisis Wavelet.*

El uso del Análisis Wavelet ha sido investigado en los últimos años, debido a su habilidad de, aproximadamente, decorrelacionar data autocorrelacionada [14]. Los coeficientes Wavelet proveen información compacta acerca de una señal en diferentes localizaciones de tiempo y frecuencia. Lo que permite también reducir la dimensionalidad de la data, pero en este caso en la dirección del tiempo a diferencia del PCA o el FDA que lo hace en la dirección de las variables. Esto trae como ventaja una disminución del tiempo de cálculo y en algunos casos reducir el error de clasificación debido a la eliminación de las componentes ruidosas de la señal. Debido a las razones antes descritas el Análisis Wavelet es utilizado en este trabajo como una de las técnicas para realizar la etapa de extracción de características del sistema de clasificación de patrones. Una descripción matemática resumida de la Transformada Discreta Wavelet (TDW) se muestra a continuación, se recomienda a los lectores interesados en una explicación detallada consultar [15].

Una señal puede ser escrita como:

$$s(t) = \sum_k c_{j_{\max}}(k) \varphi_{j_{\max},k}(t) + \sum_{j=1}^{j_{\max}} \sum_k d_j(k) \psi_{j,k}(t) \quad (2.23)$$

Donde:

$$\varphi_{j,k}(t) = 2^{-j/2} \varphi(2^{-j}t - k) \quad (2.24)$$

$$\psi_{j,k}(t) = 2^{-j/2} \psi(2^{-j}t - k) \quad (2.25)$$

son familias de funciones generadas de funciones básicas de escalamiento $\varphi(t)$ y la función wavelet $\psi(t)$, para el escalamiento y traslación respectivamente. El índice k determina la localización de la wavelet en el dominio del tiempo, mientras que el índice j determina la ubicación en el dominio de la frecuencia, lo que permite una localización tiempo-frecuencial. Los coeficientes de aproximación y de detalle vienen dados por:

$$c_j(k) = \sum_m h(m - 2k) c_{j+1}(m) \quad \text{Coeficientes de Aproximación} \quad (2.26)$$

$$d_j(k) = \sum_m g(m - 2k) c_{j+1}(m) \quad \text{Coeficientes de Detalle} \quad (2.27)$$

En muchas señales de los procesos industriales la información más importante esta contenida en las bajas frecuencias, mientras que los contenidos de altas frecuencias están relacionados con ruido. En el Análisis Wavelet los contenidos de baja frecuencia se encuentran en los coeficientes de aproximación y los de alta frecuencia en los coeficientes de detalle.

2.5.5. Análisis Discriminante Generalizado (GDA).

En el diagnóstico de fallas el Análisis Discriminante Generalizado (GDA) fue utilizado como una de las técnicas para realizar la etapa de extracción de características del sistema de clasificación de patrones.

Para el caso en que la data requiere de un clasificador no lineal, el clásico FDA ha sido generalizado a su versión kernel, en el trabajo de Mika [20] fue formulado el Discriminante de Fisher Kernel (KFD) para el caso de clasificación binaria, mientras que Baudat y Anouar [16] desarrollaron el Análisis Discriminante Generalizado (GDA) para problemas multiclases.

El uso de las funciones kernel permite realizar una proyección no lineal desde el espacio original R^N a un espacio característico F de alta dimensionalidad: $\phi: z \in R^N \rightarrow \phi(z) \in F$, donde las diferentes clases pueden separadas por un plano. Un ejemplo de un problema de clasificación de este tipo se muestra en la Figura 2.5, donde la data en el espacio original R^2 , requiere límites de decisión no lineal del tipo elipsoidal para separar las clases A y B . La data es proyectada en un espacio característico de mayor dimensión R^3 , donde un hiperplano lineal puede separar las dos clases.

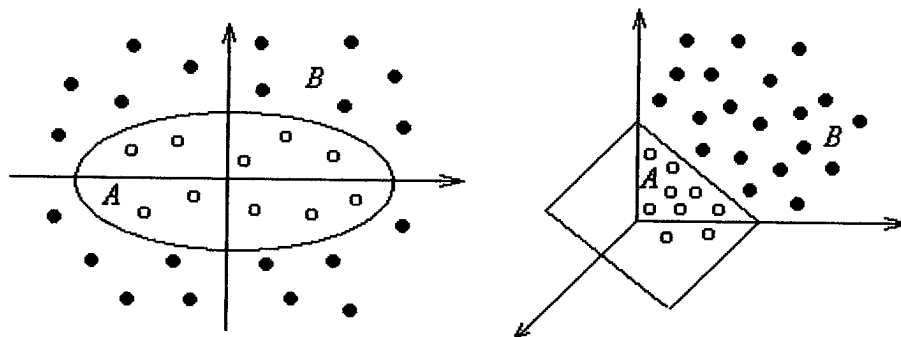
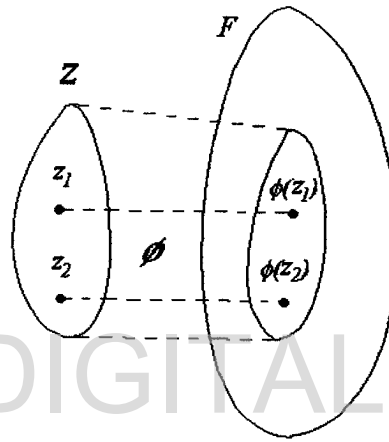


Figura 2.5. **Izquierda:** Data de las Clases A y B en el espacio original 2-D. **Derecha:** Data proyectada en un espacio característico 3-D.

El espacio característico F puede ser considerado como un “espacio de linealización”, sin embargo su dimensión puede ser arbitrariamente grande y posiblemente infinita. Afortunadamente no se necesita el $\phi(z)$ exacto y el espacio característico puede volverse implícito usando los métodos kernels, donde productos punto en F son reemplazados con una función kernel en el espacio original R^N tal que la proyección no lineal es realizada implícitamente en R^N . (Figura 2.6).



$$k(z_1, z_2) = \phi(z_1)\phi(z_2)$$

Figura 2.6. Artificio Kernel. El producto punto en el espacio característico kernel F es definido implícitamente.

En este trabajo se utilizó la función Gaussiana RBF que es una de las más usadas funciones kernels:

$$k(z_1, z_2) = \exp\left(\frac{-\|z_1 - z_2\|^2}{\sigma^2}\right), \quad \sigma^2 \in R, \quad (z_1, z_2) \in R^N \quad (2.28)$$

Luego, la idea es realizar el clásico FDA en el espacio característico F en vez del espacio original R^N . Siendo S_{WF} la matriz de dispersión-intra-clases y S_{BF} la matriz de dispersión-entre-clases en el espacio característico F , expresadas como sigue:

$$S_{WF} = \frac{1}{L} \sum_{i=1}^C \sum_{j=1}^{C_i} (\phi_{i,j} - \bar{\phi}_i)(\phi_{i,j} - \bar{\phi}_i)^T \quad (2.29)$$

$$S_{BF} = \frac{1}{L} \sum_{i=1}^C C_i (\bar{\phi}_i - \bar{\phi})(\bar{\phi}_i - \bar{\phi})^T \quad (2.30)$$

$$\phi_{i,j} = \phi(z_{ij}) \quad , \quad \bar{\phi}_i = \frac{1}{C_i} \sum_{j=1}^{C_i} \phi(z_{ij}) \quad (2.31)$$

Donde C_i es el número de elementos en Z_i , por lo que $L = \sum_{i=1}^C C_i$.

De manera similar que lo hace FDA, GDA determina un conjunto óptimo de vectores discriminantes tal que se maximice la matriz de dispersión-entre-clases mientras se minimiza la matriz de dispersión-intra-clases. La maximización puede ser alcanzada resolviendo el siguiente problema de valores propios:

$$\Psi = \arg \max_{\Psi} \frac{|\Psi^T S_{BF} \Psi|}{|\Psi^T S_{WF} \Psi|}, \quad \Psi = [\Psi_1, \dots, \Psi_M], \quad \Psi_m \in F \quad (2.32)$$

La extracción de características basado en GDA de una data de entrada z puede ser obtenida por una proyección lineal en F .

$$y_m = \Psi_m((\phi(z) - \bar{\phi})) \quad (2.33)$$

En la Figura 2.7 se muestra un ejemplo que compara el resultado de aplicar el FDA lineal y el GDA no lineal a la data Iris de Fisher, compuesta por 3 clases de flores, donde cada clase contiene 4 variables y 50 observaciones.

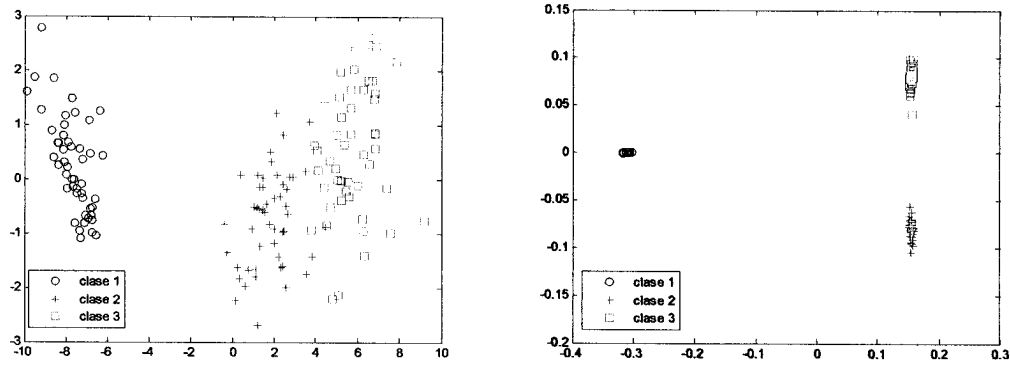


Figura 2.7. **Izquierda:** Data Iris de Fisher aplicando FDA. **Derecha:** Data Iris de Fisher aplicando GDA.

2.5.6. Análisis de pares FDA.

Una vez que una falla ha sido detectada, el siguiente paso es la identificación de la misma, realizando la Selección de las Variables (SV) que están más relacionadas con la falla ocurrida, es decir que han presentado una mayor variación. La SV permite además mejorar la clasificación en la etapa de diagnóstico como lo veremos más adelante. Para ello se utilizó el método propuesto en [21] donde las fallas son identificadas aplicando el análisis de pares FDA a la data en condiciones normales, que denotaremos como X_0 y la data en presencia de cada clase de falla, que denotaremos como X_i . La matriz de dispersión para X_0 es:

$$S_0 = \frac{1}{n_0} \sum_{x \in X_0} (x - \bar{x}_0)(x - \bar{x}_0)^T \quad (2.34)$$

Siendo $\bar{x}_0$ el vector de media de las observaciones de X_0 .

La matriz de dispersión para X_i es:

$$S_i = \frac{1}{n_i} \sum_{x \in X_i} (x - \bar{x}_i)(x - \bar{x}_i)^T \quad (2.35)$$

Siendo $\bar{x}_i$ el vector de media de las observaciones de X_i .

La matriz de dispersión-intra-clases es:

$$S_w = S_0 + S_i \quad (2.36)$$

La matriz de dispersión-entre-clases viene dada por:

$$S_b = \frac{1}{n_i} \sum_{j=0, i_i} (\bar{x}_j - \bar{x})(\bar{x}_j - \bar{x})^T \quad (2.37)$$

Siendo $\bar{x}$ el vector de media de las observaciones en X_0 y X_i ,

Al incluir sólo dos clases en el análisis de pares FDA, sustituyendo (2.36) y (2.37) en la ecuación (2.21) y resolviendo el problema de valores propios generalizado se obtendrá sólo un valor propio significativo λ_i y una dirección de Fisher, correspondiente al vector propio ϕ_i . Se define la dirección de Fisher ϕ_i como la dirección de falla para X_i . Los pesos en ϕ_i son utilizados para generar el gráfico de contribución para la falla X_i , donde:

$$\phi_i = [\phi_1, \phi_2, \dots, \phi_j, \dots, \phi_m]^T \quad (2.38)$$

El j^{th} elemento ϕ_j corresponde a la contribución de la j^{th} variable.

2.5.7. Análisis Discriminante.

Para realizar el diagnóstico de fallas, el Análisis Discriminante (AD) fue utilizado como *calculador discriminante* del sistema de clasificación de patrones. El objetivo del AD es encontrar reglas de asignación de individuos a una de las clases de una clasificación preestablecida [22].

La clasificación de las fallas se realiza asignando una observación a la clase o falla i , si esta presenta el máximo valor de la **función discriminante** [9]:

$$g_i(x) > g_j(x) \quad (2.39)$$

donde $g_i(x)$ es la función discriminante para la clase j , dado un vector de data $x \in R^m$.

Siendo w_i una condición de falla i , el error de clasificación puede ser minimizado usando la función discriminante:

$$g_i(x) = P(w_i | x) \quad (2.40)$$

donde $P(w_i | x)$ es la probabilidad posterior de que x pertenezca a la clase i . Usando la regla de Bayes:

$$P(w_i | x) = \frac{p(x | w_i)P(w_i)}{p(x)} \quad (2.41)$$

donde $P(w_i)$ es la probabilidad *a priori* para la clase w_i , $p(x)$ es la función de densidad de probabilidad para x y $p(x | w_i)$ es la función de densidad de probabilidad para x condicionado en w_i . Ha sido demostrado que una clasificación idéntica ocurre cuando (2.40) es reemplazado por:

$$g_i(x) = \ln p(x|w_i) + \ln P(w_i) \quad (2.42)$$

Asumiendo que la data para cada clase esta normalmente distribuida, $p(x|w_i)$ esta dado por:

$$p(x|w_i) = \frac{1}{(2\pi)^{m/2} [\det(\Sigma_i)]^{1/2}} \exp \left[-\frac{1}{2} (x - \bar{x}_i)^T \Sigma_i^{-1} (x - \bar{x}_i) \right] \quad (2.43)$$

Donde $\bar{x}_i$ y Σ_i son el vector de media y la matriz de covarianza para la clase i respectivamente. Sustituyendo (2.43) en (2.42) y asumiendo que la probabilidad a priori para cada clase es la misma, la **función discriminante** resulta:

$$g_i(x) = (x - \bar{x}_i)^T \Sigma_i^{-1} (x - \bar{x}_i) - \ln(\det(\Sigma_i)) \quad (2.44)$$

Si no se cuenta con suficiente data para estimar el vector de media y la matriz de covarianza de cada clase entonces la clasificación resultara en una solución subóptima, en este caso, realizar una reducción de dimensionalidad puede aplicarse para mejorar la clasificación, técnicas como el PCA y el FDA permiten realizar dicha reducción en la etapa de extracción de características.

2.6. Procedimiento propuesto para el estudio de las fallas.

A continuación se detalla el procedimiento propuesto en este trabajo para realizar todo el estudio de las fallas, que incluye los 3 principales objetivos como son la detección, identificación y diagnóstico de fallas. Este procedimiento puede ser aplicado en todos aquellos procesos industriales multivariantes en los que se cuente con una base de datos históricos de fallas ocurridas en el mismo.

En el algoritmo propuesto mostrado en la Figura 2.8 se especifican todas las etapas, que incluye para la etapa de detección el estadístico Hotelling T^2 y de ser necesario puede incluir el método DISSIM en paralelo para darle mayor robustez. Luego en la etapa de extracción de características se propone el uso del Análisis Wavelet para decorrelacionar la data si esta se encuentra autocorrelacionada, el diagnóstico de fallas se realiza en dos etapas, un primer diagnóstico por medio de FDA+AD que logra clasificar las fallas utilizando un clasificador lineal y una segunda etapa para aquellas fallas que requieren de una segunda clasificación, utilizando un clasificador no lineal, esta segunda etapa incluye un paso intermedio para la identificación de fallas por medio del Análisis de pares FDA, que permite realizar la selección de las variables (SV) más relacionadas con la falla, ya que estas serán utilizadas en la segunda clasificación realizada por medio de GDA+SV+AD.

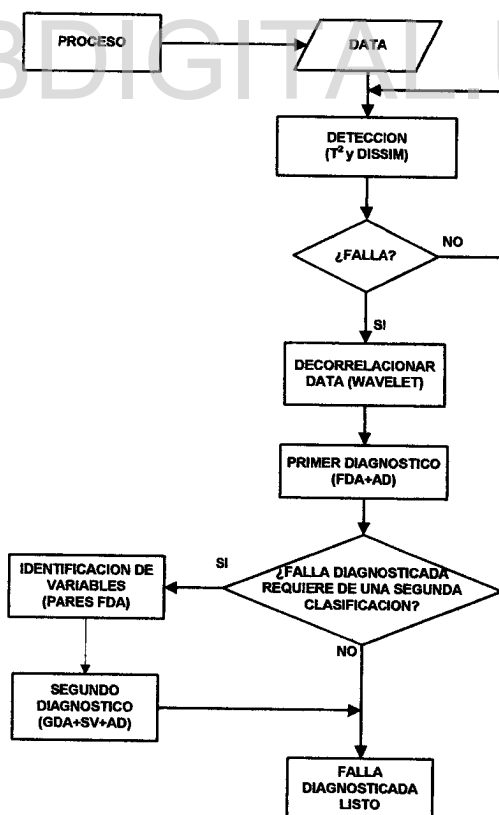


Figura 2.8. Algoritmo propuesto para la detección y diagnóstico de fallas

Capítulo III. Casos de Estudio

En este capítulo se describen los casos de estudio que permiten probar el desempeño de las técnicas estudiadas en el capítulo anterior para realizar la detección, identificación y diagnóstico de fallas. Idealmente es conveniente contar con una data histórica de un proceso real en presencia de fallas, debido a que este tipo de data no se encuentra a la mano públicamente por motivo de seguridad de las empresas, muchos investigadores han optado por probar los métodos en procesos industriales simulados en computador.

El proceso de un tanque de reacción no isotérmico agitado continuamente (CSTR) fue simulado para generar a partir de el, la data histórica de un primer caso de estudio. Este proceso se ha utilizado previamente en un buen número de publicaciones [23,24,25]. La data de la simulación de un segundo caso de estudio, el Proceso Tennessee Eastman (TEP) fue obtenida de [17], este proceso ha sido usado anteriormente como benchmark¹, para probar sistemas de detección y diagnóstico de fallas y estrategias de control, por lo que se puede conseguir en una gran cantidad de publicaciones [9,12,19,26,27,28].

Todos los resultados que se muestran en los siguientes capítulos fueron obtenidos a partir de la base de datos históricos de los dos procesos simulados mencionados anteriormente.

¹ Benchmark: Estándar o patrón mediante el cual se pueden comparar, medir o evaluar algo.

3.1 Descripción del proceso de un tanque de reacción no isotérmico agitado continuamente (CSTR)

La descripción del proceso no lineal de un tanque de reacción no isotérmico agitado continuamente (CSTR), las ecuaciones diferenciales, las fallas y las perturbaciones que comúnmente presenta dicho proceso, así como los detalles del sistema de control cascada son obtenidos de [23].

En la Figura 3.1. se muestra un esquema del proceso del CSTR con el sistema de control cascada y la instrumentación involucrada.

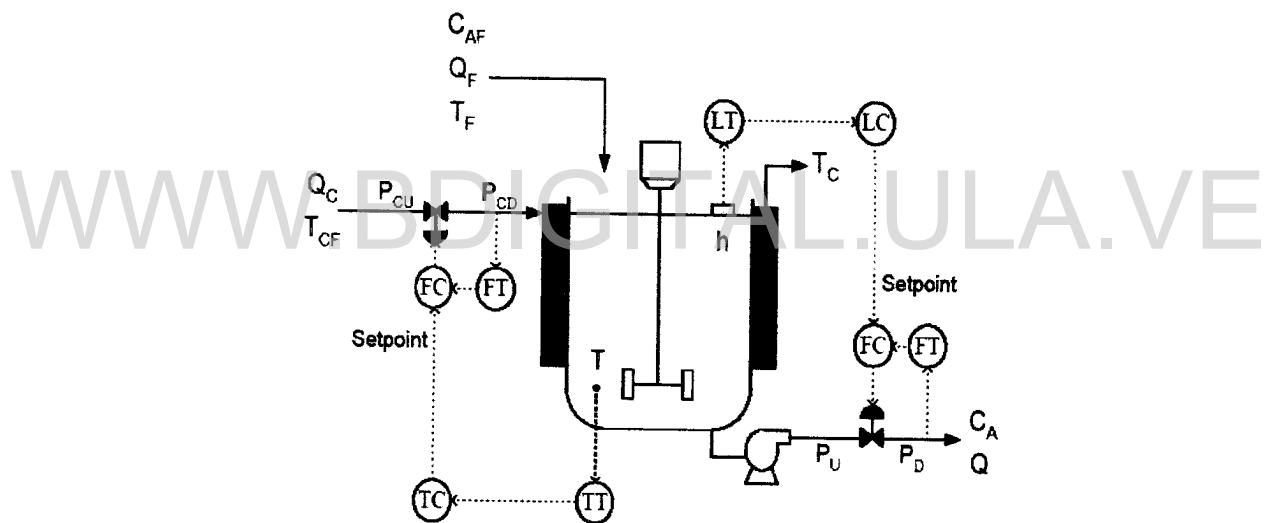


Figura 3.1. Proceso del CSTR con control cascada.

El CSTR esta compuesto por un tanque, una chaqueta enfriadora, un agitador, una bomba y la instrumentación que componen el sistema de control cascada, tales como, válvulas de control, transmisores de nivel, flujo y temperatura, así como de controladores retroalimentados del tipo PI. Las variables controladas son la temperatura del reactor (T) y el nivel del tanque (h), las variables manipuladas son el flujo de alimentación del refrigerante de la chaqueta (Q_C) y el flujo de salida del reactor (Q).

Se asume que en el interior del tanque ocurre una reacción de primer orden irreversible, $A \rightarrow B$. Un modelo dinámico del CSTR puede ser obtenido a partir de las asunciones de un mezclado perfecto y parámetros físicos constantes.

Los balances de energía y de componentes en el tanque reactor son:

$$\frac{dT}{dt} = \frac{1}{Ah} (Q_F T_F - QT) - \frac{\Delta H}{\rho C_p} C_A k_0 \exp(-E/RT) + \frac{UA_c}{\rho C_p Ah} (T_c - T) \quad (3.1)$$

$$\frac{dC_A}{dt} = \frac{1}{Ah} (Q_F C_{AF} - QC_A) - C_A k_0 \exp(-E/RT) \quad (3.2)$$

El balance de energía en la chaqueta queda:

$$\frac{dT_c}{dt} = \frac{Q_c}{V_c} (T_{CF} - T_c) + \frac{UA_c}{\rho_c C_{pc} V_c} (T - T_c) \quad (3.3)$$

El balance de masa en el tanque reactor queda:

$$\frac{dh}{dt} = \frac{Q_F - Q}{A} \quad (3.4)$$

donde:

h = nivel del reactor (dm)

Q_F = flujo de alimentación del reactor (L/min)

Q = flujo de salida del reactor (L/min)

Q_c = flujo de alimentación del refrigerante (L/min)

A = área transversal del reactor (dm²)

T = temperatura en el reactor (K)

T_F = temperatura en el flujo de alimentación del reactor (K)

ΔH = calor de reacción (J/mol)

ρ = densidad del contenido del reactor (g/L)

ρ_C = densidad del refrigerante (g/L)

C_P = capacidad calorífica del contenido del reactor (J/(g K))

C_{PC} = capacidad calorífica del refrigerante (J/(g K))

C_A = concentración de las especies A en el reactor (mol/L)

C_{AF} = concentración de las especies A en el flujo de alimentación del reactor (mol/L)

k_0 = constante de reacción (min^{-1})

E = energía de activación (J/mol)

R = constante de los gases (J/(mol K))

U = coeficiente de transferencia de calor (J/(min K dm^2))

A_C = área de transferencia de calor (dm^2)

T_C = temperatura del refrigerante en la chaqueta (K)

T_{CF} = temperatura en el flujo de alimentación del refrigerante (K)

V_C = volumen de la chaqueta (L)

Tabla 3.1. Parámetros y condiciones de operación nominal.

$Q_F = 100,02 \text{ L/min}$	$h = 6,0 \text{ dm}$
$Q = 100 \text{ L/min}$	$A = 16,66 \text{ dm}^2$
$Q_C = 15 \text{ L/min}$	$k_0 = 7,2 \times 10^{10} \text{ min}^{-1}$
$T_F = 320 \text{ K}$	$\Delta H = -5 \times 10^4 \text{ J/mol}$
$T_{CF} = 300 \text{ K}$	$\rho C_P = 239 \text{ J/(L K)}$
$T = 402,35 \text{ K}$	$\rho_C C_{PC} = 4175 \text{ J/(L K)}$
$T_C = 345,44 \text{ K}$	$E/R = 8750 \text{ K}$
$C_{AF} = 1,0 \text{ mol/L}$	$U A_C = 5 \times 10^4 \text{ J/(min K)}$
$C_A = 0.037 \text{ mol/L}$	$V_C = 10 \text{ L}$

3.1.1 Sistema de Control

El sistema de control del CSTR tiene dos variables controladas (T y h) y dos variables manipuladas (Q y Q_C). T y h son controlados a partir de Q_C y Q respectivamente.

Un sistema de control del tipo cascada es implementado con el fin de corregir el efecto de las perturbaciones en el flujo de salida del reactor y el flujo de alimentación del refrigerante antes de que afecten las variables controladas. La configuración del control cascada es mostrada en la Figura 3.1.

Los transmisores de flujo presentes en el sistema de control cascada son modelados sólo como una ganancia debido a que generalmente son rápidos y los transmisores de nivel y temperatura son modelados por medio de una función de transferencia de primero orden [29]. La instrumentación se asume analógica de 4-20 mA. Las ganancias y constantes de tiempo se muestran en la Tabla 3.2:

Tabla 3.2. Parámetros de modelado de los transmisores del sistema de control

Medición	Ganancia (K)	Constante de Tiempo (τ)
Q (L/min)	16 mA/(200 L/min))	—
Q_C (L/min)	16 mA/(30 L/min))	—
T (K)	16 mA/(200 K)	0,33 min
h (dm)	16 mA/(3 dm)	0,33 min

Las válvulas de control son modeladas por medio de una función de transferencia de primero orden,

$$G_v = \frac{K}{\tau s + 1} \quad (3.5)$$

donde $K = 1/16 \text{ mA}^{-1}$ y $\tau = 2 \text{ s}$. El flujo a través de cada válvula de control viene dado por:

$$Q = C_v f(l) \sqrt{\frac{\Delta P_v}{g_s}} \quad (3.6)$$

Q = flujo (gal/min)

C_v = coeficiente de la válvula

$f(l)$ = característica de la válvula

l = apertura de la válvula (0=cerrado, 1=completamente abierto)

ΔP_v = caída de presión a través de la válvula (psi)

g_s = gravedad específica del fluido

Se asume la característica de la válvula de tipo lineal y la caída de presión a través de cada válvula constante. Las presiones nominales en la línea de descarga del reactor a la entrada y salida de la válvula, P_E y P_S respectivamente, las presiones nominales en la línea de alimentación del refrigerante a la entrada y salida de la válvula, P_{RE} y P_{RS} respectivamente, y los coeficientes de las válvulas son mostrados en la Tabla 3.3.

Tabla 3.3. Presiones nominales y coeficientes de las válvulas.

Descarga del reactor	Alimentación del refrigerante
$P_E = 35,0 \text{ psi}$	$P_{RE} = 24,7 \text{ psi}$
$P_S = 14,7 \text{ psi}$	$P_{RS} = 14,7 \text{ psi}$
$C_v = 44,4$	$C_v = 9,5$

Controladores del tipo PI fueron usados para los controladores esclavos y maestros. Los parámetros de sintonización son mostrados en la Tabla 3.4.

Tabla 3.4. Parámetros de los controladores PI

Lazo de control	K_c	τ_I (s)
Q (L/min)	0,2	1,7
Q_c (L/min)	0,2	1,7
T (K)	-2,0	52,2
h (dm)	-3,0	90,0

3.1.2 Descripción de las Fallas.

Las 14 fallas mostradas en la Tabla 3.5 componen una gran cantidad de fallas y perturbaciones que pueden ser encontradas en una típica base de datos históricos de los proceso CSTR. Las primeras 12 fallas fueron obtenidas de [23], las fallas 13 (Backlash² en la válvula de refrigerante) y 14 (Válvula del flujo de salida del reactor atascada) fueron incluidas en este trabajo y corresponden a fallas de válvula que se presentan comúnmente. La falla 13 es representada por un juego en las partes mecánicas de la válvula, que puede ser vista como una zona muerta al cambiar la dirección de la señal de control y tiene un valor numérico igual al 20% del alcance de la válvula. La falla 14 no tiene valor numérico y es representada por un valor constante del flujo de la salida del reactor sin importar cual sea la señal del controlador.

Las fallas 4,6,7,8,9,10,11 y 12, pueden ser tanto positivas como negativas, lo que resulta en un total de 22 fallas que fueron simuladas para la generación de la base de datos históricos.

² Backlash: De acuerdo a la Sociedad Americana de Instrumentación (ISA) "En instrumentación de procesos, Backlash es un movimiento relativo entre las partes mecánicas interactuando, cuando el movimiento es invertido, resultado de la existencia de juego entre ellas.

Tabla 3.5. Fallas y perturbaciones del CSTR

FALLA	DESCRIPCION	VALOR
1. Desactivación Catalítica	La energía de activación aumenta en rampa por 100min.	La rampa para E/R es +3K/min
2. Falla en el Intercambiador de Calor	El coeficiente UAc baja en rampa por 100min.	Rampa para UAc de -125 (J/(min-K))/min
3. Falla en el transmisor de Flujo de Refrigerante	El Flujo de Refrigerante se mantiene a su último valor por 100 min.	N/A
4p y 4n. Bias en la medición de Temperatura (T)	Bias en la medición de temperatura del reactor por 100 min.	± 4 K
5. Stiction ³ en la válvula de refrigerante	La posición de la válvula de refrigerante no cambia a menos que el valor de la posición deseada difiera de la posición actual 5% o más del alcance de la válvula.	N/A
6p y 6n. Cambio escalón en Q_F	El Flujo de alimentación se mantiene en su nuevo valor por 100 min.	± 10 L/min
7p y 7n. Cambio tipo rampa en C_{AF}	La Concentración de la alimentación aumenta y disminuye en rampa por 100 min.	La rampa es de $\pm 6 \times 10^{-4}$ (mol/L)/min
8p y 8n. Cambio tipo rampa en T_F	La Temperatura de la alimentación aumenta y disminuye en rampa por 100 min.	La rampa es de $\pm 0,1$ K/min
9p y 9n. Cambio tipo rampa en T_{CF}	La Temperatura de la alimentación del Refrigerante aumenta y disminuye en rampa por 100 min.	La rampa es de $\pm 0,1$ K/min
10p y 10n. Cambio escalón en P_{RE}	La presión en la línea de alimentación del refrigerante a la entrada de la válvula se mantiene en el nuevo valor por 100 min.	$\pm 2,5$ psi
11p y 11n. Cambio escalón en P_S	La presión en la línea de descarga del reactor a la salida de la válvula se mantiene en el nuevo valor por 100min.	± 5 psi
12p y 12n. Cambio en el Set Point de Temperatura	El S.P. para la temp. del reactor se mantiene en el nuevo valor por 100 min	± 3 K
13. Backlash en la válvula de refrigerante	La válvula del refrigerante presenta juego.	20% del alcance
14. Válvula del flujo de salida del reactor atascada	La válvula del flujo de salida del reactor se mantiene a su última apertura por 100 min.	N/A

³Stiction: De acuerdo a la Sociedad Americana de Instrumentación (ISA) “stiction es la resistencia al comienzo del movimiento, usualmente medido como la diferencia entre el valor del impulso requerido para superar la fricción estática”. De acuerdo a [30] “stiction es una propiedad de un elemento tal que su movimiento suave o uniforme en respuesta a una fuerza aplicada precede de un momento en que se encuentra estático seguido por un repentino y abrupto salto llamado slip-jump. El slip-jump viene expresado como un porcentaje del alcance de la válvula. Su origen en un sistema mecánico es debido a que la fricción estática excede la fricción dinámica durante el inicio de un movimiento”.

3.1.3 Variables medidas del proceso CSTR.

Las variables del proceso CSTR (Tabla 3.6) que fueron incluidos en la base de datos históricos corresponde aquellas variables que se consideran relacionadas a las fallas y perturbaciones y que pueden ser medidas con instrumentación común. A continuación se muestran las variables seleccionadas.

Tabla 3.6. Variables medidas del Proceso CSTR

Símbolo	Descripción
C_A	Concentración del reactor
T	Temperatura del reactor
T_C	Temperatura del refrigerante
h	Nivel del reactor
Q	Flujo de salida del reactor
Q_C	Flujo del refrigerante
Q_F	Flujo de alimentación del reactor
C_{AF}	Concentración de alimentación
T_F	Temperatura de alimentación
T_{CF}	Temperatura de alimentación del refrigerante
h_C	Señal de control de nivel
Q_C	Señal de control de Flujo
T_C	Señal de control de temperatura
Q_{CC}	Señal de control del flujo del refrigerante

Para realizar la medición de las 14 variables del proceso, fue necesario incluir 6 transmisores más, a parte de los pertenecientes al sistema de control, estos corresponden a las mediciones de C_A , C_{AF} , Q_F , T_F , T_{CF} , T_C . Los transmisores de flujo son modelados sólo como una ganancia y los transmisores de concentración y temperatura son modelados por medio de una función de transferencia de primero orden. La instrumentación se asume analógica de 4-20 mA. Las ganancias y constantes de tiempo se muestran en la Tabla 3.7.

Tabla 3.7. Parámetros de modelado de los transmisores

Medición	Ganancia (K)	Constante de Tiempo (τ)
C_A (mol/L)	16 mA/(0,18 mol/L)	0,5 min
C_{AF} (mol/L)	16 mA/(0,4 mol/L)	0,5 min
Q_F (L/min)	16 mA/(40 L/min))	—
T_F (K)	16 mA/(40 K)	0,33 min
T_{CF} (K)	16 mA/(40 K)	0,33 min
T_C (K)	16 mA/(230 K)	0,33 min

3.1.4 Generación de la data histórica.

La data histórica fue generada por medio de simulación, utilizando el software Simulink del paquete Matlab, en un procesador Pentium IV de 3.0 GHz y 1.0 GB de memoria RAM. La condición normal de operación y cada una de las fallas fueron simuladas por un tiempo de 150 min introduciendo la falla luego de transcurridos 50 min de la corrida. Mediciones de 14 variables del proceso mostradas en la Tabla 3.6 fueron almacenadas cada 5 seg, el número de observaciones generadas para cada corrida fue $n=1801$ para cada una de las variables ($m=14$).

Se generaron dos conjuntos de datas para el entrenamiento y la validación. La data de entrenamiento consistió en una corrida para el caso sin falla y una corrida para cada una de las 22 fallas, no existiendo simultaneidad en la ocurrencia de la falla. La data de validación consistió en 10 corridas para la data sin falla y para cada falla se realizaron 10 corridas donde el tamaño de la falla o perturbación para los casos 3,5,13 y 14 era fijo y para el resto fue cambiado aleatoriamente en un porcentaje que variaba entre el 25% al 125%, esto con el fin de probar la capacidad de generalización del sistema de detección y diagnóstico de fallas propuesto. Se generaron en total 230 corridas para la validación, resultando un total de 414.230 puntos para cada variable.

La mayoría de las variables medidas y algunas variables del proceso fueron contaminadas con ruido blanco con el fin de simular la variabilidad presente en un proceso real, la semilla (seed) del ruido fue cambiada aleatoriamente para cada una de las corridas. Los valores de la desviación estándar del ruido son tomados del mismo trabajo de [23] y mostrados a continuación en la Tabla 3.8.

Tabla 3.8. Desviación estándar del ruido.

Variable de proceso	Desviación estándar
Q (L/min)	0,667
Q_C (L/min)	0,1
T (K)	0,667
C_A (mol/L)	$3,33 \times 10^{-5}$
h (dm)	$1,6 \times 10^{-3}$
Q_F (L/min)	0,667
C_{AF} (mol/L)	$3,33 \times 10^{-5}$
T_F (K)	0,667
T_{CF} (K)	0,667
T_C (K)	0,667
P_S (psi)	0,333
P_{RE} (psi)	0,236

En las Figuras 3.2 y 3.3, se muestra un ejemplo de la data generada para el entrenamiento. Donde se comparan la condición normal de operación y la Falla 4p (Bias en la medición de temperatura T) en las variables Temperatura del reactor (T) y Flujo del refrigerante (Q_c), que corresponden a la variable controlada y manipulada respectivamente.

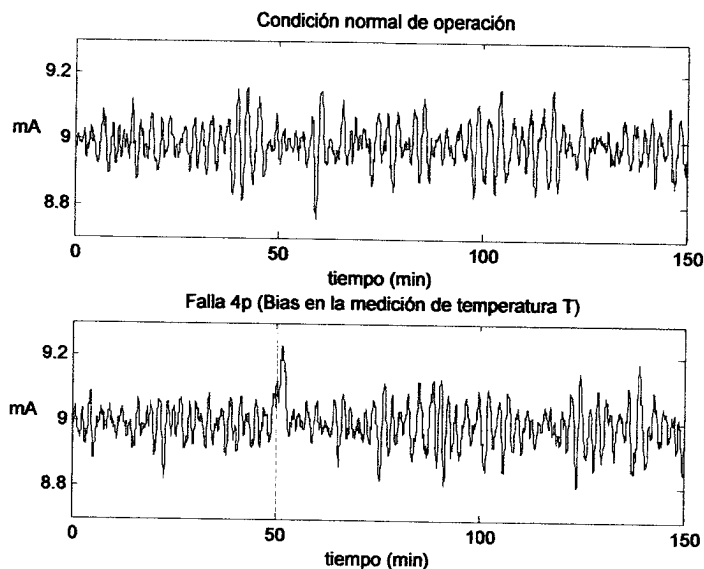


Figura 3.2. Comparación de la variable “Temperatura del reactor (T)” para la condición normal de operación y la Falla 4p (Bias en la medición de temperatura T)

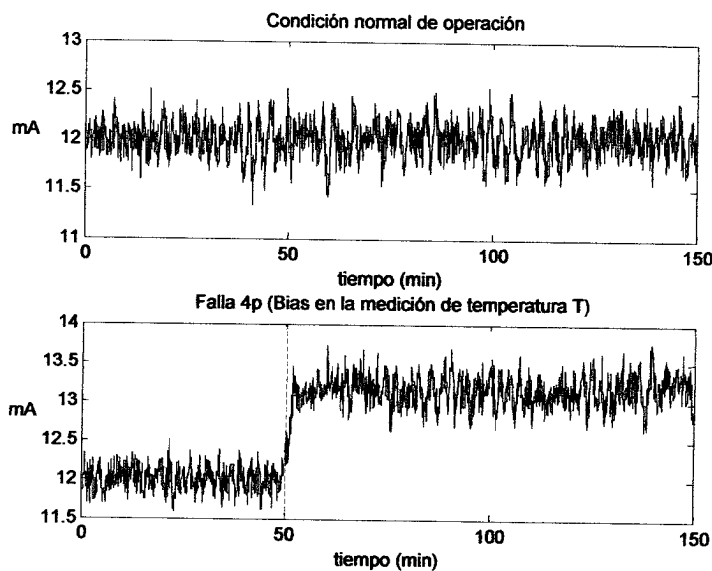


Figura 3.3. Comparación de la variable “Flujo del refrigerante (Q_c)” para la condición normal de operación y la Falla 4p (Bias en la medición de temperatura T)

3.2 Descripción del Tennessee Eastman Process (TEP)

El TEP fue creado por la compañía Eastman Chemical para proveer un proceso industrial real para evaluar técnicas de control de procesos y métodos de detección de fallas. Esta basado en una simulación de un proceso industrial real donde los componentes, dinámicas y condiciones de operación han sido modificadas por razones del propietario. El proceso consiste en 5 unidades principales, un reactor, un condensador, un compresor, un separador de vapor-liquido y una columna separadora. Que procesan 8 componentes: A, B, C, D, E, F, G y H. La descripción del proceso, la descripción de las fallas, las variables medidas, el sistema de control y otros aspectos relevantes son obtenidos de [9].

Un diagrama del TEP con el sistema de control y la instrumentación involucrada es mostrado en la Figura 3.4

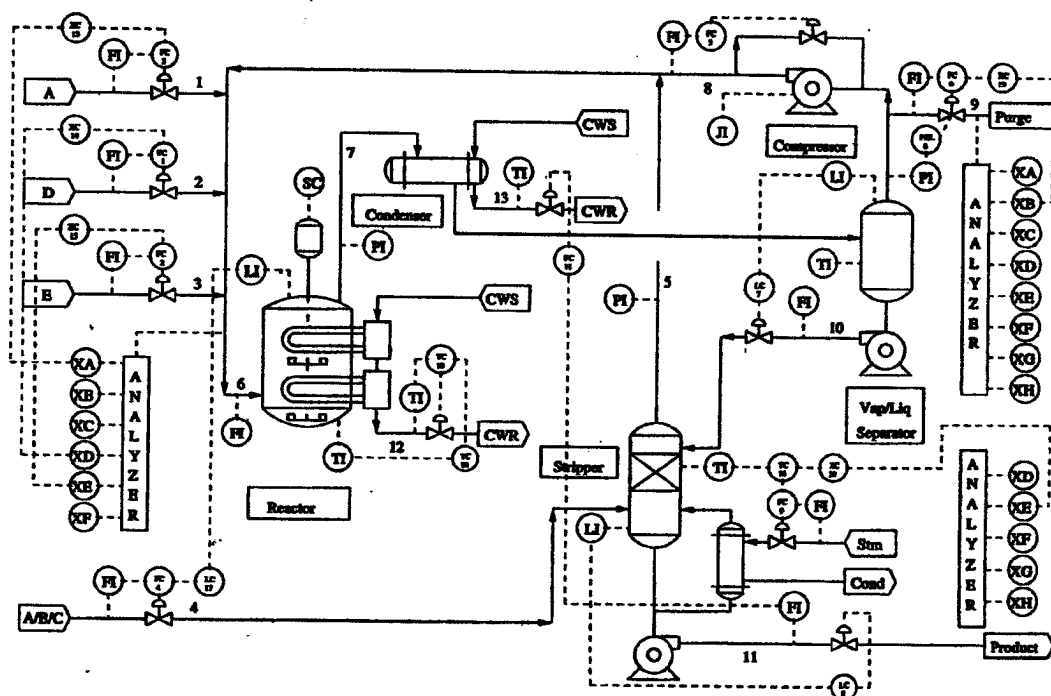
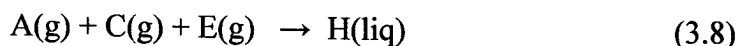
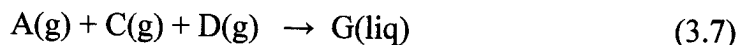


Figura 3.4. Diagrama de proceso del TEP.

Los reactantes gaseosos A, C, D, E y el inerte B son alimentados al reactor donde los productos líquidos G y H son formados. Las reacciones en el reactor son:



La especie F es un subproducto de las reacciones. Las reacciones son irreversibles, exotérmicas y aproximadamente de primer orden con respecto a las concentraciones de los reactantes. La rata de reacción son funciones de Arrhenius de temperatura, donde la reacción para G tiene una energía de activación mayor que la reacción para H, resultando en una mayor sensibilidad a la temperatura.

El flujo de los productos del reactor es enfriado en un condensador y luego alimentados a un separado de vapor-liquido. El vapor que sale del separador es reciclado a la alimentación del reactor a través de un compresor. Una porción del flujo de reciclaje es purgado para que no se acumulen el inerte y el subproducto en el proceso. Los componentes condensados del separador (Flujo 10) son bombeados a la columna separadora. El Flujo 4 es usado para separar los reactantes restantes del Flujo 10, los cuales son combinados con el flujo de reciclaje a través del Flujo 5. Los productos G y H que salen de la base de la columna separadora son enviados a un proceso aguas abajo que no esta incluido en el diagrama.

3.2.1 Sistema de Control

El sistema de control consiste de diecinueve lazos (Figura 3.4), que utilizan controladores del tipo Proporcional (P) y Proporcional e Integral (PI), la gran mayoría de ellos se encuentran en una configuración del tipo cascada.

3.2.2 Descripción de las Fallas.

En el proceso Tennessee Eastman son simuladas un total de 21 fallas (Tabla 3.9). Dieciséis de las fallas son conocidas y 5 desconocidas. Las fallas de la 1 a la 7 están relacionadas con un cambio escalón en variables del proceso. Las fallas de la 8 a la 12 están relacionadas con un incremento en la variabilidad de algunas variables del proceso. La falla 13 es un cambio gradual suave en la dinámica de la reacción y por último las Fallas 14,15 y 21 están relacionadas con válvulas que presentan stiction.

Tabla 3.9. Fallas del TEP

Falla	Descripción	Tipo
1	Relación del flujo de alimentación de A/C, composición de B constante (corriente 4)	Escalón
2	Composición de B, relación del flujo de alimentación de A/C constante (corriente 4)	Escalón
3	Temperatura del flujo de alimentación de D (corriente 2)	Escalón
4	Temperatura del agua de enfriamiento a la entrada del reactor	Escalón
5	Temperatura del agua de enfriamiento a la entrada del condensador	Escalón
6	Fuga en el flujo de alimentación de A (corriente 1)	Escalón
7	Perdida de presión en la línea de C - Disponibilidad reducida (corriente 4)	Escalón
8	Composición del flujo de Alimentación de A,B,C (corriente 4)	Variacion Aleatoria
9	Temperatura del flujo de alimentación de D (corriente 2)	Variacion Aleatoria
10	Temperatura del flujo de alimentación de C (corriente 4)	Variacion Aleatoria
11	Temperatura del agua de enfriamiento a la entrada del reactor	Variacion Aleatoria
12	Temperatura del agua de enfriamiento a la entrada del condensador	Variacion Aleatoria
13	Dinamicas de la reacción	Cambio Gradual Suave
14	Valvula del agua de enfriamiento del reactor	Stiction
15	Valvula del agua de enfriamiento del condensador	Stiction
16	Desconocida	
17	Desconocida	
18	Desconocida	
19	Desconocida	
20	Desconocida	
21	La valvula de la corriente 4 se mantiene en una posición fija	Posición Constante

3.2.3 Variables del TEP.

El proceso contiene 41 variables medidas mostradas en la Tabla 3.10 y 12 variables manipuladas mostradas en la Tabla 3.11. De las variables medidas, 22 de ellas

corresponden a mediciones del proceso y las 19 restantes corresponden a mediciones de composición realizadas en los flujos 6,9 y 11 para los diferentes componentes.

Tabla 3.10. Variables medidas del TEP

Variable	Descripción	Unidades
Vmed(1)	Flujo Alimentación de A (corriente 1)	kscmh
Vmed(2)	Flujo Alimentación de D (corriente 2)	kg/h
Vmed(3)	Flujo Alimentación de E (corriente 3)	kg/h
Vmed(4)	Flujo Alimentación total (corriente 4)	kscmh
Vmed(5)	Flujo de reciclaje (corriente 8)	kscmh
Vmed(6)	Flujo de alimentación del reactor (corriente 6)	kscmh
Vmed(7)	Presión del reactor	kPa
Vmed(8)	Nivel del reactor	%
Vmed(9)	Temperatura del reactor	°C
Vmed(10)	Flujo de purga (corriente 9)	kscmh
Vmed(11)	Temperatura del producto en el separador	°C
Vmed(12)	Nivel del producto en el separador	%
Vmed(13)	Presión del producto en el separador	kPa
Vmed(14)	Flujo a la salida inferior del separador (corriente 10)	m ³ /h
Vmed(15)	Nivel de la columna separadora	%
Vmed(16)	Presión de la columna separadora	kPa
Vmed(17)	Flujo a la salida inferior de la columna separadora (corriente 11)	m ³ /h
Vmed(18)	Temperatura de la columna separadora	°C
Vmed(19)	Flujo de vapor de la columna separadora	kg/h
Vmed(20)	Potencia del compresor	kW
Vmed(21)	Temperatura del agua de enfriamiento a la salida del reactor	°C
Vmed(22)	Temperatura del agua de enfriamiento a la salida del separador	°C
Vmed(23)	Componente A (corriente 6)	% molar
Vmed(24)	Componente B (corriente 6)	% molar
Vmed(25)	Componente C (corriente 6)	% molar
Vmed(26)	Componente D (corriente 6)	% molar
Vmed(27)	Componente E (corriente 6)	% molar
Vmed(28)	Componente F (corriente 6)	% molar
Vmed(29)	Componente A (corriente 9)	% molar
Vmed(30)	Componente B (corriente 9)	% molar
Vmed(31)	Componente C (corriente 9)	% molar
Vmed(32)	Componente D (corriente 9)	% molar
Vmed(33)	Componente E (corriente 9)	% molar
Vmed(34)	Componente F (corriente 9)	% molar
Vmed(35)	Componente G (corriente 9)	% molar
Vmed(36)	Componente H (corriente 9)	% molar
Vmed(37)	Componente D (corriente 11)	% molar
Vmed(38)	Componente E (corriente 11)	% molar
Vmed(39)	Componente F (corriente 11)	% molar
Vmed(40)	Componente G (corriente 11)	% molar
Vmed(41)	Componente H (corriente 11)	% molar

Tabla 3.11. Variables manipuladas del TEP

Variable	Descripción
Vman(1)	Flujo alimentación de D (corriente 2)
Vman(2)	Flujo alimentación de E (corriente 3)
Vman(3)	Flujo alimentación de A (corriente 1)
Vman(4)	Flujo alimentación total (corriente 4)
Vman(5)	Flujo de reciclaje del compresor
Vman(6)	Flujo de purga (corriente 9)
Vman(7)	Flujo del separador (corriente 10)
Vman(8)	Flujo de la columna separadora (corriente 11)
Vman(9)	Flujo de vapor de la columna separadora
Vman(10)	Flujo del agua de enfriamiento del reactor
Vman(11)	Flujo del agua de enfriamiento del condensador
Vman(12)	Velocidad del agitador

3.2.4 Generación de la data histórica.

La data histórica fue obtenida directamente de [18]. Los conjuntos de data de entrenamiento y de validación incluyen todas las variables medidas y las variables manipuladas, excepto la velocidad del agitador del reactor para un total de 52 variables observadas ($m=52$). El tiempo de muestreo es de 3 min. Todas las variables del proceso incluyen ruido blanco.

La data en el conjunto de entrenamiento consiste de 22 corridas, donde la semilla (seed) del ruido fue cambiada aleatoriamente para cada una de las corridas, una corrida corresponde al caso sin falla y las restante 21 corridas corresponden a cada una de las fallas mostradas anteriormente en la Tabla 3.9. El tiempo de la simulación para cada corrida fue de 25 horas. La simulación comienza sin falla y luego de transcurrida 1 hora de simulación se introduce la falla. Sólo fueron recolectados los datos después de introducida la falla lo que generó un total de 480 observaciones ($n=480$).

La data en el conjunto de validación consiste de 22 corridas, donde la semilla (seed) del ruido fue cambiada aleatoriamente para cada una de las corridas. El tiempo de la simulación para cada corrida fue de 48 horas. La simulación comienza sin falla y luego de transcurridas 8 horas de simulación se introduce la falla. Se recolectaron los datos de las 48 horas de simulación lo que generó un total de 960 observaciones ($n=960$).

En las Figuras 3.5 y 3.6, se muestra un ejemplo de la data generada para la validación. Donde se comparan la condición normal de operación y la Falla 1 “Cambio escalón en la relación del flujo de alimentación de A/C, composición de B constante (corriente 4)” en las variables “Componente A (corriente 6) (% molar)” y “Flujo Alimentación de A (corriente 1)”, que corresponden a la variable controlada y manipulada respectivamente.

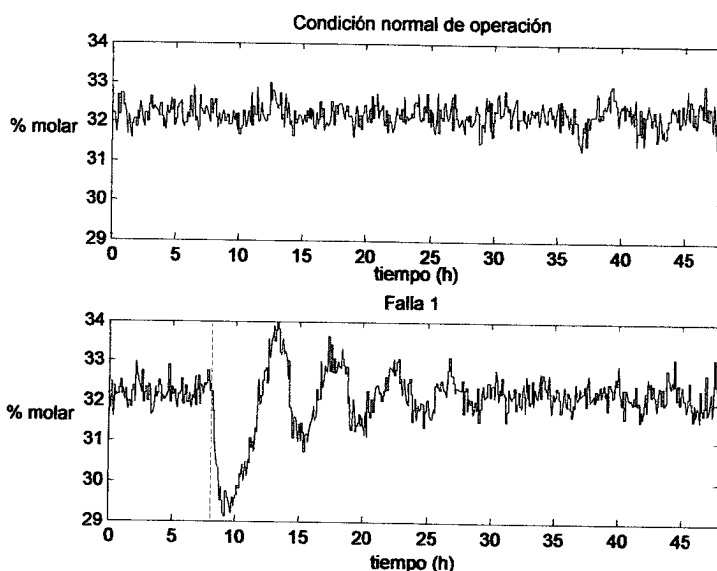


Figura 3.5. Comparación de la variable “Componente A (corriente 6) (% molar)” para la condición normal de operación y la Falla 1 (Cambio escalón en la relación del flujo de alimentación de A/C, composición de B constante (corriente 4))

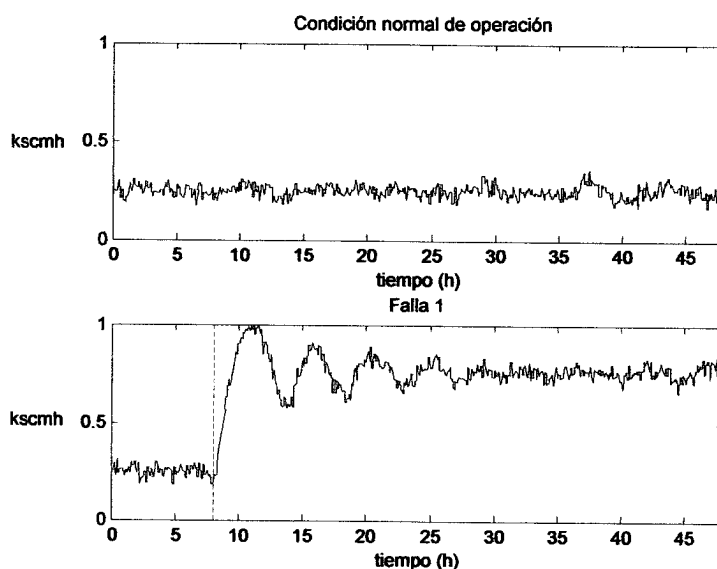


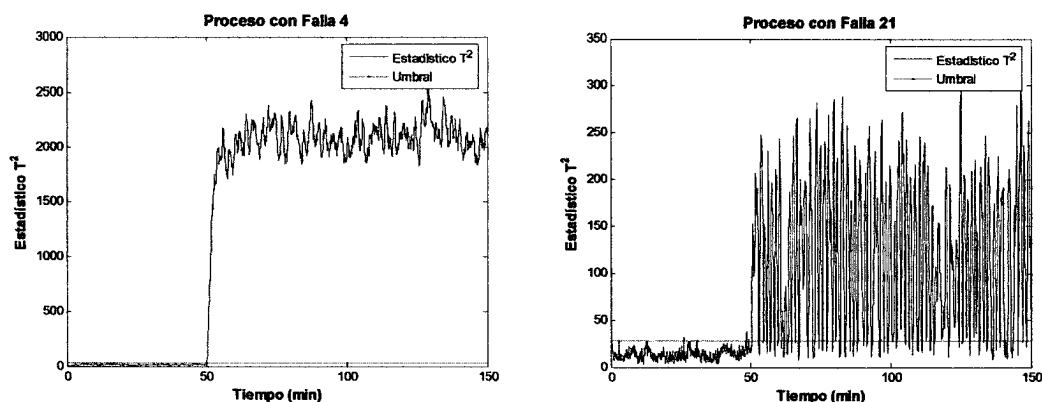
Figura 3.6. Comparación de la variable “Flujo Alimentación de A (corriente 1)” para la condición normal de operación y la Falla 1 (Cambio escalón en la relación del flujo de alimentación de A/C, composición de B constante (corriente 4)).

Capítulo IV. Resultados y Discusión.

En este capítulo se muestran los resultados de aplicar los métodos de detección y diagnóstico de fallas descritos en el capítulo 2, en el proceso de un tanque de reacción no isotérmico agitado continuamente (CSTR) y en el proceso Tennessee Eastman (TEP) escogidos como casos de estudios y detallados en el capítulo 3. Los métodos de detección de las fallas son evaluados por medio de la confiabilidad que es una medida del número de puntos por encima del umbral en el tiempo que transcurre la falla, siendo 100% el valor ideal y por el tiempo de retardo en detectar la falla, siendo un valor de 0 una respuesta inmediata. Los métodos de clasificación de las fallas son evaluados por medio del error de clasificación %ec (porcentaje de corridas que fueron clasificadas incorrectamente) y el error de clasificación individual %eci (porcentaje de las observaciones de una corrida que fueron clasificados incorrectamente). Los resultados obtenidos en este capítulo justifican la propuesta del algoritmo para la detección y diagnóstico de fallas mostrado en la Figura 2.8.

4.1 Detección de Fallas del CSTR utilizando el estadístico Hotelling (T^2)

El estadístico Hotelling (T^2) presentó un buen desempeño en la detección de las fallas presentes en el proceso CSTR, las fallas fueron detectadas en su totalidad y la gran mayoría de ellas de manera rápida. En la Figura 4.1 podemos observar la gráfica de control de T^2 ante la presencia de algunas de las fallas.


 Figura 4.1 Gráfica de control T^2 para las fallas 4 y 21.

La Tabla 4.1 muestra los tiempos de retardo en detectar la falla y la confiabilidad del estadístico T^2 para la data de entrenamiento y validación.

 Tabla 4.1. Tiempos de retardo y confiabilidad del T^2 para la data de entrenamiento y validación.

Falla	Data de Entrenamiento			Data de Validación		
	Confiabilidad (%)	Unidades de Tiempo	Tiempo (min)	Confiabilidad (%)	Unidades de Tiempo	Tiempo (min)
NF	0,99	****	****	0,98	****	****
1	93,7	86	7,2	91,7	91,3	7,6
2	91,8	104	8,7	88,5	134,6	11,2
3	99,8	5	0,4	99,4	8,6	0,7
4=4p	99,9	2	0,2	99,7	4,9	0,4
5=4n	99,9	2	0,2	99,8	4,9	0,4
6=5	63,3	5	0,4	63,5	10,2	0,9
7=6p	100,0	0	0,0	100,0	0	0,0
8=6n	100,0	0	0,0	100,0	0	0,0
9=7p	99,7	5	0,4	99,7	5,1	0,4
10=7n	99,7	5	0,4	99,6	6	0,5
11=8p	92,7	89	7,4	89,0	118,6	9,9
12=8n	93,5	81	6,8	87,0	143,9	12,0
13=9p	92,3	103	8,6	89,1	122,7	10,2
14=9n	92,7	96	8,0	89,4	125,7	10,5
15=10p	100,0	0	0,0	98,0	1,7	0,1
16=10n	100,0	0	0,0	99,8	1,2	0,1
17=11p	100,0	0	0,0	100,0	0	0,0
18=11n	100,0	0	0,0	100,0	0,3	0,0
19=12p	100,0	0	0,0	100,0	2	0,2
20=12n	100,0	0	0,0	100,0	2,6	0,2
21=13	87,4	2	0,2	85,2	7,3	0,6
22=14	99,8	5	0,4	99,5	7,8	0,7
Promedio	95,7	26,8	2,2	94,5	36,3	3,0

Licencia Creative Commons:

Atribución - No Comercial - Compartir Igual 3.0 Venezuela
(CC BY-NC-SA 3.0 VE)

4.2 Identificación de fallas del CSTR utilizando las direcciones de falla en pares FDA.

La identificación de la falla nos permite determinar las variables más relacionadas con la falla ocurrida, a partir de estas, es posible realizar una deducción de cual puede ser el tipo de falla, esto por supuesto con la ayuda de la experiencia de los operadores de la planta. La identificación de las variables también presenta una gran utilidad a la hora de diagnosticar la falla, debido a que se puede incluir sólo estas variables al aplicar los métodos de diagnóstico, lo que mejora el error de clasificación como se mostrara más adelante. En la Figura 4.2 se muestra un ejemplo de la gráfica de contribución FDA para la fallas 10 (cambio tipo rampa negativo en C_{AF}) y 12 (cambio tipo rampa negativo en T_F). Como observamos la variable que tiene mayor contribución en el gráfico del Proceso con Falla 10 es la 8, que corresponde precisamente a la Concentración de Alimentación (C_{AF}) que es la que esta siendo modificada por la falla. En el gráfico del Proceso con Falla 12 observamos que las variables que tienen una mayor contribución son la 1 y la 9, lo que desde el punto de vista del proceso es correcto ya que la variable 9 es precisamente la temperatura T_F que esta siendo modificada por la falla y la variable 1 es la concentración del reactor C_A , que es afectada directamente por la variación de T_F .

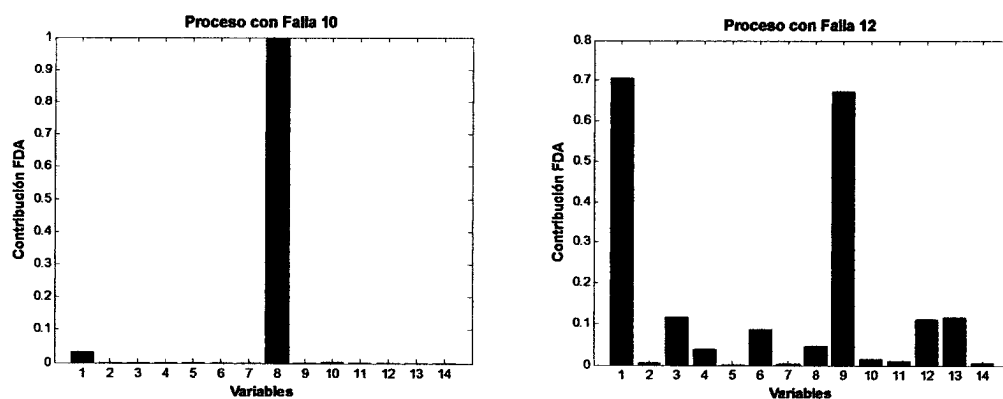


Figura 4.2 Gráfica de contribución FDA para las fallas 10 y12, data de entrenamiento.

La Tabla 4.2 muestra cuales fueron las variables identificadas para la data de entrenamiento, seleccionando las que presentaron una contribución mayor de 0,3.

Tabla 4.2 Variables Identificadas para la data de entrenamiento.

Data de Validación (td-td+100)	
Falla	Variables Identificadas
1	1
2	1,3
3	3,1,13,12
4	3,12,13,6,1
5	3,12,13,6,1
6	1,12
7	11,1,7,14,12
8	1,11,7,14
9	8
10	8
11	1,9
12	9,1
13	1,10
14	1,10
15	1,12
16	1,12
17	1,14
18	1,14
19	3,12,13,6,1
20	3,1,12,13,6
21	1,12,8
22	14,11

4.3 Clasificación de las Fallas del CSTR utilizando Análisis Discriminante (AD).

Para la clasificación de las fallas se utilizó inicialmente la técnica del Análisis Discriminante (AD) sin aplicar ningún método para realizar la etapa de extracción de características. Es importante destacar que para clasificar una falla es necesario que un patrón de esta se encuentre previamente en la data histórica, que corresponde a la data de entrenamiento. Para clasificar correctamente una falla se requiere contar con la mayor cantidad de datos que se pueda, dependiendo estos del periodo de tiempo después de detectada la falla que sea seleccionado para realizar la clasificación, un periodo de tiempo muy pequeño trae consigo una clasificación más rápida pero un

error de clasificación mayor, al contrario de un periodo de tiempo prolongado. Es deseable que la clasificación de la falla se realice en un período de tiempo corto, para evitar que el proceso no se vuelva inestable e inseguro debido al efecto prolongado de la falla en el proceso. En este trabajo el error de clasificación para las 22 fallas fue realizado para 3 periodos de tiempo después de detectada la falla por medio del estadístico Hotelling (T^2), el primero corresponde a 100 observaciones luego de detectada la falla (td-td+100), correspondiendo td al tiempo que fue detectada la falla, el cual varia dependiendo de la falla como vimos en la Tabla 4.1. El segundo período corresponde a 500 observaciones luego de detectada la falla (td-td+500) y el último corresponde a todas las observaciones después de detectada la falla (td-1801). La Tabla 4.3 muestra el error de clasificación %ec y el error de clasificación individual %eci para las 22 fallas con la data de entrenamiento y de validación utilizando el AD. La data de entrenamiento consiste en una corrida para cada caso. La data de validación consiste en 10 corridas para cada caso.

Tabla 4.3. Error de clasificación utilizando AD para diferentes periodos de tiempo con la data de entrenamiento y validación.

Falla	AD entren(td-td+100)		AD entren(td-td+500)		AD entren(td-1801)		AD valid(td-td+100)		AD valid(td-td+500)		AD valid(td-1801)	
	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci
1	0	0,0	0	0,0	0	0,0	0	0,0	0	0,0	0	0,0
2	0	0,0	0	0,0	0	0,0	0	1,7	0	0,4	0	0,2
3	0	0,0	0	0,0	0	0,0	100	100,0	100	100,0	100	100,0
4	0	1,9	0	0,4	0	0,2	0	5,8	0	1,5	0	1,0
5	0	1,9	0	0,4	0	0,2	0	10,9	0	3,9	0	2,7
6	0	20,8	0	15,0	0	15,7	0	27,4	0	24,5	0	22,8
7	0	0,0	0	0,0	0	0,0	0	4,0	0	0,8	0	0,3
8	0	0,0	0	0,0	0	0,0	0	2,1	0	0,4	0	0,2
9	0	30,2	0	8,3	0	3,5	50	48,5	0	15,8	0	6,9
10	0	0,0	0	0,4	0	0,2	0	8,3	0	6,2	0	2,9
11	0	5,7	0	2,0	0	0,9	0	19,1	0	5,5	0	2,6
12	0	0,0	0	0,0	0	0,0	0	8,1	0	2,7	0	1,3
13	0	3,8	0	1,2	0	0,5	0	14,3	0	4,2	0	2,1
14	0	3,8	0	0,8	0	0,4	0	3,6	0	1,3	0	0,6
15	0	5,7	0	2,4	0	1,0	40	47,4	40	43,2	40	42,5
16	0	7,5	0	1,6	0	0,7	40	50,4	40	44,8	40	44,1
17	0	0,0	0	0,0	0	0,0	10	11,9	10	11,8	10	11,6
18	0	1,9	0	0,4	0	0,2	30	34,7	30	33,3	30	32,7
19	0	13,2	0	2,8	0	1,2	30	43,8	30	32,9	30	31,2
20	0	11,3	0	2,4	0	1,0	30	31,3	10	17,6	10	15,7
21	0	5,7	0	7,5	0	9,3	0	22,1	0	21,2	0	20,7
22	0	5,7	0	1,2	0	0,5	0	6,2	0	1,3	0	0,6
Promedio	0,0	5,4	0,0	2,1	0,0	1,6	15,0	22,8	11,8	17,0	11,8	15,6

Podemos observar de los resultados de la Tabla 4.3 como el error de clasificación individual %eci disminuye de 22,8% para el periodo de tiempo (td-td+100) a un valor de 17,0% cuando el periodo de tiempo es (td-td+500) y disminuye aun más a 15,6% cuando el periodo de tiempo es el máximo (td-1801), lo que demuestra que el error de clasificación disminuye al contar con una mayor cantidad de datos que aporten más información pero trae consigo un retardo significativo en la clasificación.

4.4 Clasificación de las Fallas del CSTR utilizando FDA y AD.

Los resultados mostrados a continuación se obtuvieron de aplicar FDA en la etapa de extracción de características y AD para la clasificación de las fallas. Para aplicar el FDA es necesario determinar primero el número de vectores FDA o autovectores N, los resultados de la validación cruzada se muestran en la Figura 4.3. Resultando un N=12 el número de vectores FDA seleccionado.

N	FDA+AD entren(td-1801)		FDA+AD valid(td-1801)	
	% ec	%eci	% ec	%eci
1	68,1818	69,0	74,5	76,9
2	50	65,8	57,3	69,6
3	22,7273	36,1	44,1	51,4
4	0	12,6	15,9	28,1
5	0	7,4	11,8	20,6
6	0	5,8	10,0	18,7
7	0	5,1	9,1	17,4
8	0	4,5	9,1	15,9
9	0	4,4	9,1	15,9
10	0	3,6	10,5	16,1
11	0	3,3	9,1	15,2
12	0	2,3	7,7	12,4
13	0	1,9	10,0	14,2

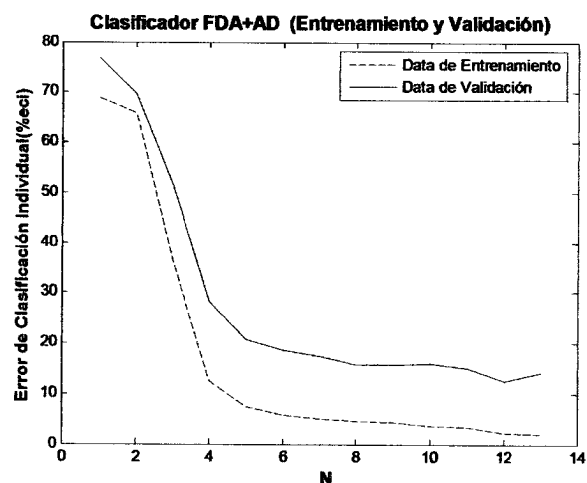


Figura 4.3. Selección del número de autovectores N por medio de validación cruzada.

La Tabla 4.4 muestra el error de clasificación %ec y el error de clasificación individual %eci para las 22 fallas con la data de entrenamiento y de validación utilizando la técnica FDA en conjunto con el AD. Para la data de validación el %ec determina el porcentaje de corridas que fueron clasificadas incorrectamente de un total de 10 corridas para cada falla y el %eci determina el porcentaje de observaciones de cada una de las 10 corridas que fueron clasificados incorrectamente.

Tabla 4.4. Error de clasificación utilizando FDA+AD para diferentes periodos de tiempo con la data de entrenamiento y validación.

N=12	FDA+AD entren(td-td+100)		FDA+AD entren(td-td+500)		FDA+AD entren(td-1801)		FDA+AD valid(td-td+100)		FDA+AD valid(td-td+500)		FDA+AD valid(td-1801)		
	Falla	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci	% ec	%eci
1	0	0	0,0	0	0,0	0	0,0	0	0,1	0	0,0	0	0,0
2	0	0	0,0	0	0,0	0	0,0	0	5,4	0	1,1	0	0,5
3	0	0	6,9	0	1,4	0	0,6	0	23,0	0	4,6	0	1,9
4	0	0	3,0	0	0,6	0	0,3	0	8,6	0	5,3	0	5,1
5	0	0	3,0	0	0,6	0	0,3	10	16,4	10	14,6	10	12,6
6	0	0	27,7	0	17,4	0	17,4	0	25,3	0	24,2	0	22,1
7	0	0	1,0	0	0,2	0	0,1	0	5,0	0	1,1	0	0,5
8	0	0	1,0	0	0,2	0	0,1	0	2,5	0	0,5	0	0,2
9	0	0	35,6	0	7,6	0	3,2	50	47,5	0	16,1	0	6,8
10	0	0	5,9	0	1,4	0	0,6	0	11,4	0	8,6	0	4,1
11	0	0	1,0	0	0,2	0	0,1	0	13,5	0	4,0	0	1,9
12	0	0	1,0	0	1,2	0	0,5	0	7,5	0	2,7	0	1,6
13	0	0	4,0	0	0,8	0	0,4	0	22,5	0	7,5	0	3,8
14	0	0	7,9	0	1,6	0	0,7	0	5,9	0	2,5	0	1,4
15	0	0	11,9	0	3,0	0	1,5	40	42,3	40	42,2	40	42,1
16	0	0	2,0	0	1,0	0	0,7	30	42,8	40	44,4	40	45,3
17	0	0	1,0	0	0,2	0	0,1	10	13,2	10	12,7	10	12,4
18	0	0	1,0	0	0,2	0	0,1	30	35,7	30	35,4	30	35,0
19	0	0	14,9	0	3,0	0	1,2	30	42,1	30	31,5	30	29,8
20	0	0	6,9	0	1,4	0	0,6	20	29,2	10	19,0	10	18,1
21	0	0	27,7	0	20,0	0	21,3	0	28,4	0	28,0	0	26,9
22	0	0	5,9	0	1,2	0	0,5	0	5,9	0	1,2	0	0,5
Promedio	0,0	0,0	7,7	0,0	2,9	0,0	2,3	10,0	19,7	7,7	14,0	7,7	12,4

Podemos observar de los resultados de las Tablas 4.3 y 4.4 como al incluir FDA en la etapa de extracción de características disminuye el %ec de 15,0% a un valor de 10,0% para la data de validación en el periodo de tiempo (td-td+100) y de 11,8% a un valor de 7,7% cuando el periodo de tiempo es (td-td+500) y (td-1801). Lo que justifica el uso del FDA como etapa previa a la clasificación.

4.5 Clasificación de las Fallas del CSTR utilizando Wavelet, FDA y AD.

Las señales antes de ser clasificadas pueden ser reducidas temporalmente, es decir reducir el número de muestras en la dirección del tiempo, utilizando el Análisis Wavelet en la etapa de extracción de características, esto trae como ventaja una disminución del tiempo de cálculo y en algunos casos reducir el error de clasificación debido a la eliminación de las componentes ruidosas de la señal. Luego de un estudio exhaustivo, por medio de validación cruzada como se muestra en la Tabla 4.5, obteniéndose el error de clasificación menor, cuando utilizamos una Wavelet del tipo Daubechies 3 (db3) y un nivel de descomposición igual a 1 y utilizando solamente los coeficientes de aproximación (ca) para formar el vector característico.

Tabla 4.5. Error de clasificación utilizando Wavelet+FDA+AD para el periodo de tiempo correspondiente a (td-1801) con la data de validación.

Nivel de Desc.	haar		db2		db3		db4		db5	
	entren %eci	valid %eci	entren %eci	valid %eci	entren %eci	valid %eci	entren %eci	valid %eci	entren %eci	valid %eci
1	2,11	12,72	2,17	11,91	2,35	11,80	2,40	11,83	2,43	11,91
2	1,72	14,78	1,92	13,59	2,16	14,59	2,10	14,33	2,05	13,91
3	1,71	16,24	1,72	16,65	2,00	15,91	1,35	15,98	1,69	16,35
4	1,59	19,76	2,11	19,21	1,83	19,44	2,01	20,21	1,57	20,27
5	2,16	26,22	3,59	24,52	3,78	24,32	4,02	26,06	4,24	26,72

La Tabla 4.6 muestra el error de clasificación %ec y el error de clasificación individual %eci para las 22 fallas con la data de validación utilizando las técnicas Análisis Wavelet y FDA en la etapa de extracción de características y el AD en la etapa de clasificación. El uso del análisis Wavelet trae como ventaja una reducción del error de clasificación promedio de 10% a 9,1% para el caso (td-td+100) y de 7,7% a 7,3% cuando el periodo de tiempo es (td-td+500) y (td-1801). Además de obtenerse una reducción de 48% de los datos, siendo este último valor muy importante para disminuir el tiempo de cálculo requerido para realizar la clasificación.

Tabla 4.6. Error de clasificación utilizando Wavelet+FDA+AD para diferentes periodos de tiempo con la data de validación.

N=12	Wavelet+FDA+AD valid(td-td+100)		Wavelet+FDA+AD valid(td-td+500)		Wavelet+FDA+AD valid(td-1801)		Falla Diag.
	% ec	%eci	% ec	%eci	% ec	%eci	
1	0	0,0	0	0,0	0	0,0	1
2	0	5,3	0	1,1	0	0,5	2
3	0	37,4	0	8,0	0	3,5	3
4	0	7,5	0	2,3	0	1,7	4
5	0	11,7	0	4,7	0	3,7	5
6	0	26,8	0	25,9	0	23,4	6
7	0	3,8	0	0,9	0	0,6	7
8	0	1,9	0	0,4	0	0,2	8
9	40	47,2	0	16,3	0	7,0	9-10
10	0	12,5	0	9,1	0	4,3	10
11	0	17,0	0	4,6	0	2,2	11
12	0	7,7	0	3,2	0	1,8	12
13	0	22,5	0	7,9	0	4,1	13
14	0	4,7	0	2,2	0	1,2	14
15	40	43,6	40	42,9	40	42,9	15-6-21
16	40	44,7	40	44,4	40	45,2	16-21
17	10	13,2	10	14,3	10	14,1	17-22
18	30	34,9	30	35,3	30	35,0	18-22
19	30	41,3	30	30,2	30	27,6	19-2
20	10	27,5	10	15,5	10	13,7	20-2
21	0	28,1	0	27,2	0	26,4	21
22	0	5,5	0	1,1	0	0,5	22
Promedio	9,1	20,2	7,3	13,5	7,3	11,8	

4.6 Clasificación de las Fallas del CSTR utilizando Wavelet, GDA y AD.

En la Tabla 4.6 se observa que las fallas 9, 15-20, fueron clasificadas incorrectamente en varios de los casos, lo que hace suponer que la data no puede ser separada por medio de un clasificador lineal como el FDA en conjunto con el AD y justifica el uso del GDA en la etapa de extracción de características. Debido al gasto computacional que requiere el método GDA, no es factible realizar la clasificación de las fallas utilizando todas las posibles fallas en la data de entrenamiento, por lo que sólo se clasificó con GDA las fallas que presentaron un error de clasificación en el análisis FDA, formándose subgrupos con aquellas fallas que se han diagnosticado incorrectamente. En la Tabla 4.6 se observa una columna que corresponde a la falla

diagnosticada, de allí se formaron los subgrupos de fallas. En la Tabla 4.7 se muestran los 4 subgrupos que han sido formados. Se incluyo además el uso de la selección de variables (SV) para mejorar la clasificación en ambas técnicas FDA y GDA, esta selección de variables se realizó utilizando las variables identificadas por medio de las direcciones de fallas en pares FDA para la data de validación, seleccionando sólo aquellas variables que tenían una mayor contribución.

Tabla 4.7. Error de clasificación de subgrupos de fallas, utilizando SV, Wavelet, FDA, GDA y AD, para el periodo de tiempo correspondiente a (td-td+100) con la data de validación.

Subgrupo	Falla	Wavelet+FDA+AD		SV+Wavelet+FDA+AD		Wavelet+GDA+AD		SV+Wavelet+GDA+AD		SV
		% ec	N	% ec	N	% ec	σ^2	% ec	σ^2	
1	9	0	3	0	1	10	46	0	1	8
	10	10		0		0		0		
2	6	0		0		0		0		
	15	30	12	30	2	30	71	30	0,5	1,12
	16	30		30		30		30		
	21	0		0		0		0		
3	17	0		0		0		0		
	18	0	2	0	1	0	76	0	6	1,4,14
	22	0		0		0		0		
4	2	0		0		0		0		
	19	20	6	20	2	10	50	0	1	1,3
	20	10		10		10		0		
Promedio		8,3		7,5		7,5		5,0		

Como se observa en la Tabla 4.7 el uso del GDA reduce el error de clasificación en comparación con el FDA, de un valor de 8,3% a 7,5% sin selección de variables y de 7,5% a 5,0% con selección de variables. Observándose también el beneficio de realizar previamente la SV e incluir sólo esas variables en la clasificación. Al realizar la clasificación de las fallas en dos etapas, un primer diagnóstico por medio de Wavelet+ FDA+AD que logra clasificar las fallas linealmente separables y una segunda etapa por medio de SV+Wavelet+GDA+AD para aquellas fallas requieren de una segunda clasificación se obtiene un error promedio para las 22 fallas de 2,7%, valor que satisface las expectativas del sistema de detección y diagnóstico de fallas propuesto.

4.7 Detección de Fallas del TEP utilizando el estadístico Hotelling (T^2)

La Tabla 4.8 muestra los tiempos de retardo y la confiabilidad del estadístico en detectar las fallas de la data de validación. Las fallas 3 y 9 no fueron detectadas y la confiabilidad del estadístico promedio de todas las fallas fue de un 80,7 %. El tiempo de retardo promedio no incluye las fallas 3 y 9. En la Figura 4.4 podemos observar en las gráficas de control de T^2 como el estadístico no detecta las fallas 3 y 9 a diferencia de las fallas 2 y 8 por ejemplo que si son detectadas de manera rápida y confiable.

Tabla 4.8. Tiempos de retardo y confiabilidad del T^2 para la data de validación.

Falla	Data de Validación		
	Confiabilidad (%)	Unidades de Tiempo	Tiempo (min)
NF	2,4	****	****
1	99,8	5	15
2	98,9	16	48
3	5,4	ND	ND
4	99,8	2	6
5	100,0	3	9
6	100,0	2	6
7	100,0	2	6
8	98,0	23	69
9	4,9	ND	ND
10	89,4	27	81
11	73,4	13	39
12	99,9	4	12
13	95,5	47	141
14	100,0	3	9
15	12,1	581	1743
16	92,4	12	36
17	95,6	26	78
18	90,6	86	258
19	86,9	13	39
20	90,4	73	219
21	61,3	263	789
Promedio	80,7		189,6

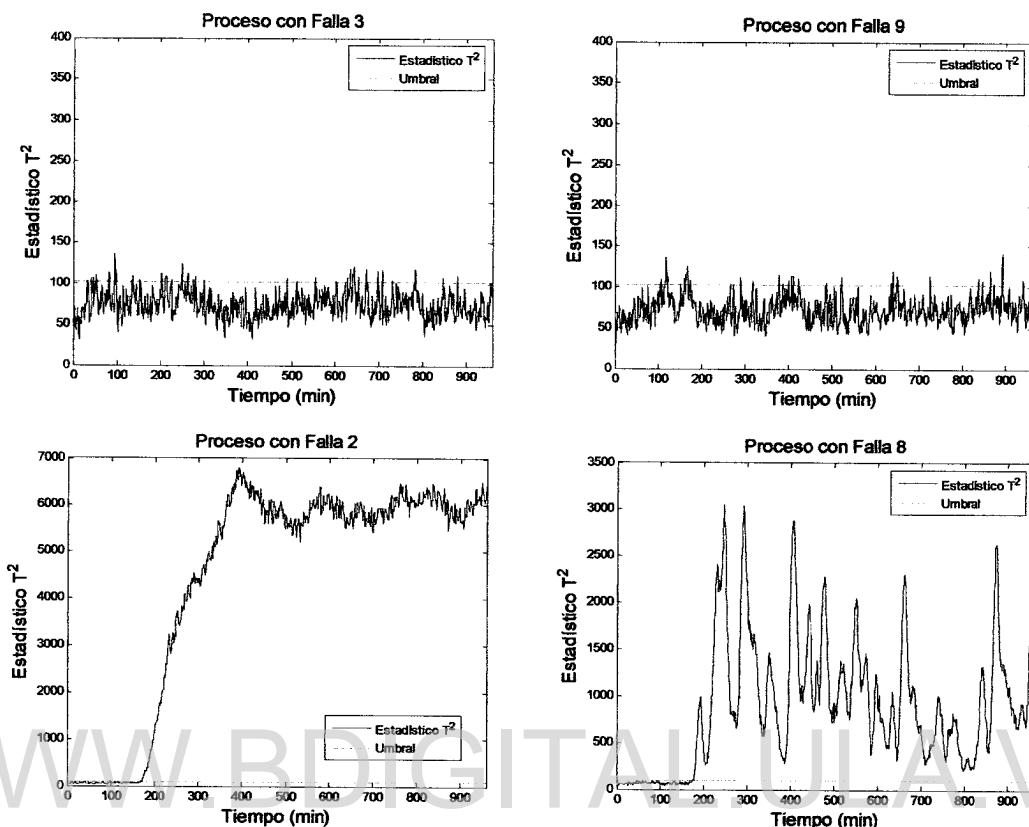


Figura 4.4 Gráfica de control de T^2 para la falla 2,3, 8 y 9.

4.8 Detección de Fallas del TEP utilizando el método DISSIM.

Debido a que el estadístico Hotelling (T^2) no fue capaz de detectar alguna de las fallas y el valor de la confiabilidad de 80,7% no representa una garantía para la detección de las fallas, se justifica entonces el uso de un método más complejo como lo es el DISSIM. La Tabla 4.9 muestra como se mejora la confiabilidad promedio para la data validación de 80,7% a un valor de 88%, detectándose además las fallas 3 y 9 que no eran detectadas por el estadístico T^2 .

Tabla 4.9. Tiempos de retardo y confiabilidad del DISSIM para la data de validación.

Data de Validación			
Falla	Confiabilidad (%)	Unidades de Tiempo	Tiempo (min)
NF	1,8	****	****
1	99,0	9	27
2	97,9	18	54
3	56,6	1	3
4	98,8	11	33
5	99,0	9	27
6	99,9	2	6
7	98,3	15	45
8	96,4	30	90
9	35,4	99	297
10	93,3	1	3
11	99,8	3	9
12	100,0	1	3
13	90,0	81	243
14	99,4	6	18
15	20,4	635	1905
16	100,0	1	3
17	100,0	1	3
18	89,0	89	267
19	96,3	31	93
20	91,4	69	207
21	88,5	9	27
Promedio	88,0		160,1

En la Figura 4.5 se observan las gráficas de control de DISSIM para el caso de las fallas 2 y 9. La línea punteada representa el tiempo donde se introduce la falla en el proceso.

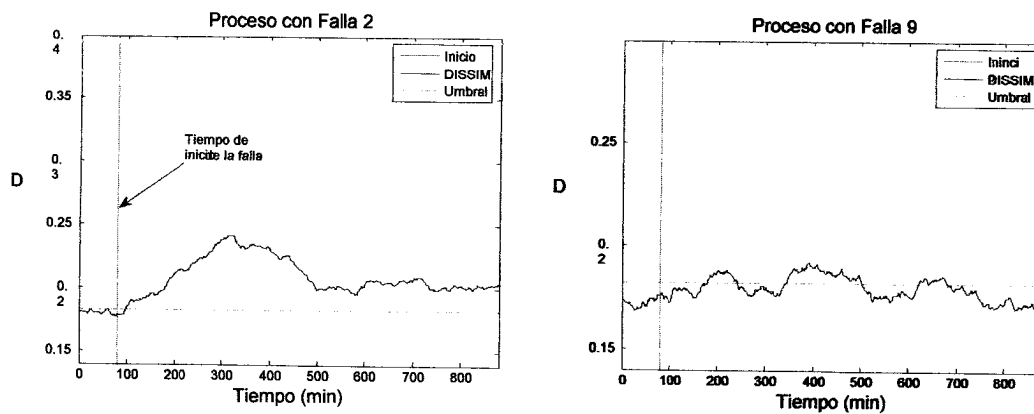


Figura 4.5 Gráfica de control del DISSIM para las fallas 2 y 9.

De los resultados obtenidos del estadístico T^2 y del DISSIM se puede observar que el estadístico T^2 a pesar de no ser capaz de detectar algunas fallas que si fueron detectadas por el DISSIM, lo realizó de manera más rápida que este en 12 de las fallas que si fueron detectadas, por lo tanto más que un método sustituir al otro se complementan. Por esta razón se propone en este trabajo el uso de ambas técnicas en paralelo para dar mayor robustez al sistema de detección de fallas.

4.9 Identificación de fallas del TEP utilizando las direcciones de falla en pares FDA.

En la Figura 4.6 se muestra un ejemplo de la gráfica de contribución FDA para las fallas 4 (Cambio escalón en la temperatura del agua de enfriamiento a la entrada del reactor) y 6 (Fuga en el flujo de alimentación de A (corriente 1)). Como observamos la variable que tiene mayor contribución en el gráfico del Proceso con Falla 4 es la 51, que corresponde precisamente al “Flujo del agua de enfriamiento del reactor” que es la variable manipulada para contrarrestar el efecto de la falla por medio del lazo de control 10 (Figura 3.4). En el gráfico del Proceso con Falla 6 observamos que las variables que tienen una mayor contribución son la 16 y la 44, esta relación es correcta desde el punto de vista del proceso ya que la variable 16 corresponde a la “Presión de la columna separadora” que se ve modificada al existir una fuga en el flujo en la corriente 1 ya que esta se encuentra conectada con la corriente 5 que es donde se encuentra el transmisor de presión (Figura 3.4) y la variable 44 corresponde precisamente al “Flujo alimentación de A (corriente 1)” que es la variable manipulada para contrarrestar el efecto de la falla por medio del lazo de control 3.

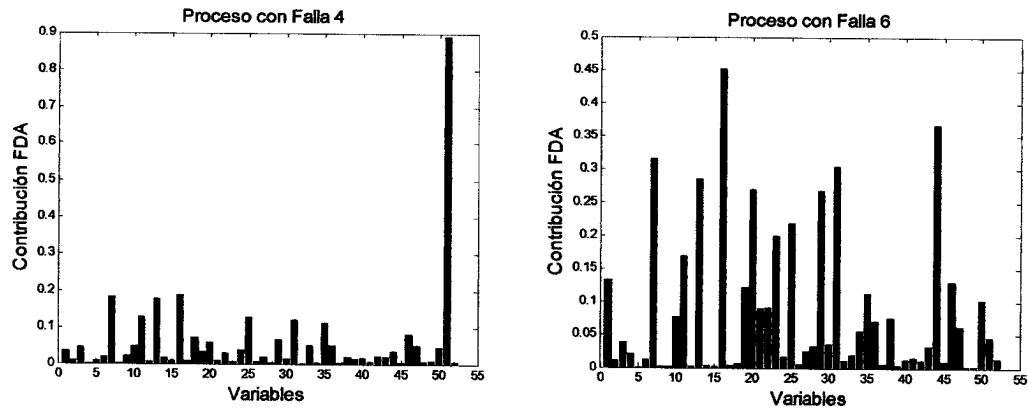


Figura 4.6 Gráfica de contribución FDA para las fallas 4 y 6, data de validación y (td-td+100).

La Tabla 4.10 muestra cuales fueron las variables identificadas para la data de entrenamiento, seleccionando únicamente aquellas que presentaron una contribución mayor de 0,3.

Tabla 4.10 Variables Identificadas para la data de validación.

Data de Validación (td-td+100)	
Falla	Variables Identificadas
1	1,44,50
2	10,47
3	7,16,13
4	51
5	19,50
6	16,44,7,31
7	45
8	10,47,16
9	50,19,18,34
10	19,50
11	51
12	19,38,50
13	19,50,16,7,18,13
14	40
15	18,50,19,31
16	50
17	21
18	22
19	35,31
20	46,20,13
21	19,50,18

4.10 Clasificación de las Fallas del TEP utilizando Análisis Discriminante (AD).

Para la clasificación de las fallas se utilizó inicialmente la técnica el AD sin aplicar ningún método para realizar la etapa de extracción de características. La clasificación para las 21 fallas fue realizado para 3 periodos de tiempo después de detectada la falla. Seleccionado el tiempo de retardo menor entre el estadístico Hotelling (T^2) y el DISSIM. El primer periodo corresponde a 100 observaciones luego de detectada la falla (td-td+100), correspondiendo td al tiempo que fue detectada la falla, el cual varia dependiendo de la falla como vimos en las Tablas 4.8 y 4.9. El segundo período corresponde a 500 observaciones luego de detectada la falla (td-td+500) y el último corresponde a todas las observaciones después de detectada la falla (td-800). La Tabla 4.11 muestra el error de clasificación %ec y el error de clasificación individual %eci para las 21 fallas con la data de entrenamiento y de validación utilizando el AD.

Tabla 4.11. Error de clasificación utilizando AD con la data de validación.

Falla	AD valid(td-td+100)		AD valid(td-td+500)		AD valid(td-800)	
	% ec	%eci	% ec	%eci	% ec	%eci
1	0	15,8	0	3,2	0	2,0
2	0	1,0	0	0,2	0	0,1
3	0	80,2	100	74,5	100	78,0
4	0	12,9	0	19,8	0	17,5
5	0	13,9	0	2,8	0	1,8
6	0	0	0	0	0	0
7	0	0	0	0	0	0
8	0	0	0	1,6	0	1,0
9	100	77,2	100	76,8	100	75,4
10	0	23,8	0	14,6	0	13,1
11	0	34,7	0	24,4	0	24,3
12	0	1,0	0	0,6	0	1,8
13	0	7,9	0	20,6	0	20,6
14	0	2,0	0	0,6	0	1,3
15	0	72,3	100	85,0	100	85,0
16	0	18,8	0	17,6	0	19,3
17	0	27,7	0	16,6	0	15,0
18	0	30,7	0	10,0	0	23,5
19	0	1,0	0	2,6	0	3,8
20	0	3,0	0	5,4	0	5,1
21	0	0	0	0	0	0,1
Promedio	4,8	20,2	14,3	17,9	14,3	18,5

4.11 Clasificación de las Fallas del TEP utilizando FDA y AD.

Para realizar la clasificación de las fallas utilizando FDA es necesario determinar primero el número de vectores FDA o autovectores N , los resultados de la validación cruzada se muestran en la Figura 4.7, resultando un $N=51$ el número de vectores FDA seleccionado.

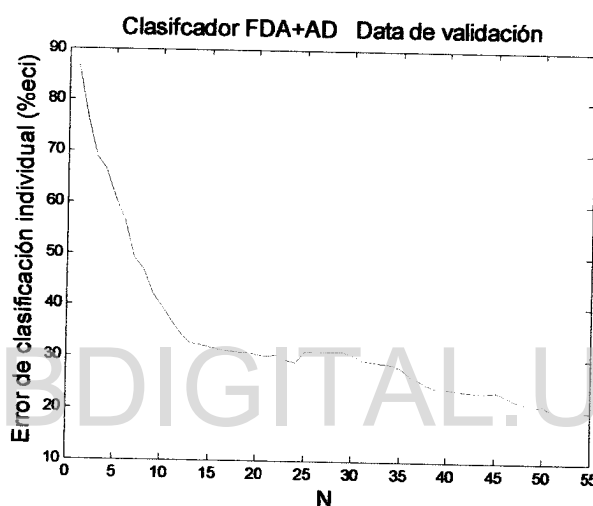


Figura 4.7 Selección del número de autovectores N por medio de validación cruzada.

La Tabla 4.12 muestra el error de clasificación %ec y el error de clasificación individual %eci para las 21 fallas con la data de validación utilizando la técnica FDA en conjunto con el AD. Se observa en los resultados de las Tablas 4.11 y 4.12 como el %ec para el periodo de tiempo $(td-td+100)$ no disminuye al aplicar FDA en la etapa de extracción, manteniéndose en el mismo valor de 4,8%. Cuando el periodo de tiempo es $(td-td+500)$ y $(td-800)$ el %ec si disminuye de 14,3% a 9,5%. Lo que justifica el uso del FDA como etapa previa a la clasificación.

Tabla 4.12. Error de clasificación utilizando FDA+AD con la data de validación.

N=51	FDA valid(td-td+100)		FDA valid(td-td+500)		FDA valid(td-800)		Falla Diag.
	% ec	%eci	% ec	%eci	% ec	%eci	
1	0	14,9	0	3,0	0	1,9	1
2	0	1,0	0	0,2	0	0,1	2
3	0	77,2	0	72,9	0	73,6	3
4	0	15,8	0	19,8	0	18,8	4
5	0	15,8	0	3,2	0	2,0	5
6	0	0	0	0	0	0	6
7	0	0	0	0	0	0	7
8	0	0	0	1,2	0	0,8	8
9	100	81,2	100	78,6	100	76,6	9-15
10	0	26,7	0	16,6	0	14,3	10
11	0	25,7	0	20,8	0	22,4	11
12	0	4,0	0	1,2	0	2,6	12
13	0	9,9	0	21,6	0	21,6	13
14	0	2,0	0	1,0	0	1,5	14
15	0	77,2	100	86,4	100	86,4	15-9
16	0	18,8	0	17,6	0	19,8	16
17	0	30,7	0	17,4	0	16,1	17
18	0	29,7	0	9,4	0	6,6	18
19	0	1,0	0	2,6	0	3,6	19
20	0	5,9	0	11,6	0	11,7	20
21	0	20,8	0	17,2	0	18,9	21
Promedio	4,8	21,8	9,5	19,1	9,5	19,0	

4.12 Clasificación de las Fallas del TEP utilizando Wavelet, FDA y AD.

Luego de un estudio exhaustivo por medio de validación cruzada como se muestra en la Tabla 4.13, se observó que para este caso en particular no se consigue mejorar el error de clasificación utilizando el análisis Wavelet, por el contrario se aumento de 19,0 %eci a 20,0 %eci en el mejor de los casos, lo que indica que esta data no se encuentra significativamente auto-correlacionada y al aplicar el análisis Wavelet se remueven información importante de la señal.

Tabla 4.13. Error de clasificación utilizando Wavelet+FDA+AD para el periodo de tiempo correspondiente a (td-800) con la data de validación.

Nivel de Desc.	harr %eci	db2 %eci	db3 %eci	db4 %eci	db5 %eci
1	20,00	23,25	24,18	26,35	25,05
2	25,63	25,43	24,99	25,37	28,41

4.13 Clasificación de las Fallas del TEP utilizando Wavelet, GDA y AD.

En la Tabla 4.12 se observa que las fallas 9 y 15, fueron clasificadas incorrectamente, lo que hace suponer que la data no puede ser separada por medio de un clasificador lineal como el FDA en conjunto con el AD y justifica el uso del GDA en la etapa de extracción de características para mejorar el error de clasificación. Para aplicar el método GDA es conveniente hacerlo sólo en el subgrupo de fallas que están presentando el error de clasificación debido al gasto computacional como se explicó anteriormente. En este caso se formó un solo subgrupo que incluye las fallas que se están confundiendo que son la 9 y la 15.

Los resultados de aplicar el método GDA en comparación con el FDA se muestran en la Tabla 4.14. Observamos como el método GDA reduce el %eci de 48,5% a 42,1% sin selección de variables y a 36,1% con selección de variables y la falla es clasificada correctamente como se observa en el valor de 0% error de clasificación (%ec). Se observa además que el uso de selección de variables (SV) mejora notablemente la clasificación en la técnica GDA, esta selección de variables se realizó previamente por medio de las direcciones de falla en pares FDA, seleccionándose las variables 18, 19, 31, 34 y 50 que eran las que tenían una mayor contribución a la falla.

Tabla 4.14. Error de clasificación del subgrupo de fallas, utilizando SV, FDA, GDA y AD, para el periodo de tiempo correspondiente a (td-td+100) con la data de validación.

FDA+AD					SV+FDA+AD				GDA+AD			SV+GDA+AD			
Subgrupo	Falla	% ec	% eci	N	% ec	% eci	N	SV	% ec	% eci	σ^2	% ec	% eci	σ^2	SV
1	9	0	46,5	32	100	65,3	5	18,19,31,	0	44,6	11	0	37,6	0,56	18,19,31,
	15	100	50,5		0	31,7		34,50	0	39,6		0	34,6		34,50
Promedio		50,0	48,5		50,0	48,5			0,0	42,1		0,0	36,1		

WWW.BDIGITAL.ULA.VE

Capítulo V. Conclusiones

En este trabajo se propuso un nuevo procedimiento para la detección, identificación y diagnóstico de fallas, que incluye el uso del estadístico Hotelling (T^2) y DISSIM para la detección de fallas, el Análisis Wavelet para decorrelacionar la data autocorrelacionada, el uso del Análisis Discriminante de Fisher (FDA) en conjunto con el Análisis Discriminante (AD) como una primera clasificación, y el uso del Análisis Discriminante Generalizado (GDA) en conjunto con el AD, utilizando selección de variables, para la clasificación de aquellas fallas donde el FDA+AD no presenta buen rendimiento, obteniéndose como resultado final todo un algoritmo que permite realizar el estudio de fallas de manera completa. Los resultados obtenidos en la detección y diagnóstico de fallas de procesos tales como el de un tanque de reacción no isotérmico agitado continuamente (CSTR) y el Tennessee Eastman Process (TEP), demostraron la eficiencia del algoritmo propuesto al obtenerse errores de clasificación muy bajos.

Los resultados obtenidos mostraron que el GDA puede ser aplicado para realizar el diagnóstico de fallas, en aquellos casos donde se requiere de un clasificador no lineal, y técnicas de clasificación lineales tales como el FDA no presentan buenos resultados.

Una limitante que se observó del uso del GDA es el gasto computacional elevado para clasificar muchas clases de fallas, lo que obligó a utilizarlo en subgrupos más pequeños. Esto es debido al uso de las funciones kernel que realizan una proyección no lineal desde el espacio original R^N a un espacio característico F de alta dimensionalidad, que para el caso del proceso CSTR cuando se realizaba el

entrenamiento para clasificar las 22 clases de fallas, este espacio de alta dimensionalidad podía llegar a ser de 13266 observaciones x 13266 variables, lo que requiere un uso de memoria RAM y un tiempo de cálculo demasiado alto, no viable para fines prácticos. Por lo tanto se utilizó en subgrupos más pequeños de 2, 3 o 4 clases de fallas.

Técnicas como el estadístico Hotelling (T^2) y el método DISSIM pueden ser utilizadas en paralelo para aumentar la robustez y la velocidad de los sistemas de detección de fallas, esto se demostró de los resultados obtenidos con el estadístico T^2 que a pesar de no ser capaz de detectar algunas fallas que si fueron detectadas por el DISSIM, lo realizó de manera más rápida en buena parte de los casos.

El uso del Análisis Wavelet en la etapa de extracción de características del sistema de clasificación de patrones permite decorrelacionar data autocorrelacionada y reducir la dimensionalidad de la misma, trayendo como ventaja una disminución del tiempo de cálculo y en algunos casos reducir el error de clasificación debido a la eliminación de las componentes ruidosas de la señal. En muchos procesos industriales esto es realmente útil, debido a que las variables son usualmente medidas a diferentes tiempos de muestreos, en ocasiones mayores de lo necesario, en otros casos con data faltante y con ruido.

El uso de selección de variables (SV) por medio del análisis de pares FDA, mejora la clasificación al aplicar las técnicas FDA y GDA, al ser seleccionadas sólo aquellas variables que tienen una mayor contribución con la falla, obteniéndose errores de clasificación menores.

REFERENCIAS BIBLIOGRÁFICAS

- [1] Venkatasubramanian, V., Rengaswamy, R., Kavuri, S., Yin K. *A review of process fault detection and diagnosis*. Computers and Chemical Engineering. 27. 293-346. (2003).
- [2] Gertler, J. *Fault Detection and Diagnosis in Engineering Systems*. Mercel Dekker. New York. (1999).
- [3] Willsky A. *A survey of design methods for failure detection in dynamics systems*. Automatica. Vol 12. (1976).
- [4] Garces, Sayra. *Propuesta de un Sistema Automático para Detección y Diagnóstico de Fallas*. Proyecto de Grado para optar al título de Magíster Scientae en Automatización e Instrumentación. Universidad de los Andes. Mérida. (1999).
- [5] Sanchez, Betsy. *Detección y Diagnóstico de Fallas utilizando Estructuras de Transición*. Proyecto de Grado para optar al título de Magíster Scientae en Automatización e Instrumentación. Universidad de los Andes. Mérida. (2002).
- [6] Zhou, Y., Hahn, J., Mannan, M. *Fault detection and classification in chemical processes based on neural networks with feature extraction*. ISA Transactions 42. 651-664. (2003).

[7] Rios, A., Rivas, F. y Moussalli, G. *Un método basado en redes neuronales para el diagnóstico y reconstrucción de fallas*. IV Congreso de Automatización y Control (CAC 2003). Mérida-Venezuela.

[8] Padilla, Delfina. *Aplicación de métodos estadísticos multivariantes en la predicción y diagnóstico de fallas en motores de encendido por compresión*. Proyecto de Grado para optar al título de Magíster Scientae en Estadística Aplicada. Universidad de los Andes. Mérida. (1997).

[9] Chiang, L., Russell, E. and Braatz, R. *Fault detection and diagnosis in industrial systems*. Springer-Verlag. Great Britain. (2001).

[10] Guillen, Marcos. *Detección y Diagnóstico de Fallas utilizando la Transformada Wavelet*. Proyecto de Grado para optar al título de Magíster Scientae en Automatización e Instrumentación. Universidad de los Andes. Mérida. (2004).

[11] Ruiz, Diego. *Fault diagnosis in chemical plants integrated to the information system*. Thesis for the degree of Doctor of the Universitat Politècnica de Catalunya, España. (2001).

[12] Zhou, Y. *Data driven process monitoring based on neural networks and classification trees*. Thesis for the degree of Doctor of Philosophy. Texas A&M University, USA. (2004).

[13] Kano, M., Hasebe, S., Hashimoto, I. and Ohno, H. *A new multivariate statistical process monitoring method using principal component analysis*. Comput. Chem. Eng. 25. 1103–1113. (2001).

- [14] Aradhye, H. B., Bakshi, B. R., Strauss, R. A., and Davis, J. F. *Multiscale SPC using wavelets: Theoretical analysis and properties*. AIChE Journal, 49, 4, 939-958 (2003).
- [15] Burrus, C., Gopinath, C. and Guo, H. *Introduction to Wavelets and Wavelet Transform*. Prentice Hall. (1998).
- [16] Baudat, G. and Anouar, F. *Generalized discriminant analysis using a kernel approach*. Neural Computation, vol. 12, 2385–2404. (2000).
- [17] Koscor, A. and Tóth, L. *Kernel-Based Feature Extraction with a Speech Technology Application*. IEEE Transactions on Signal Processing, vol. 52, No. 8, 2250-2263. (2004).
- [18] Braatz R. *Large Scale Systems Research Laboratory*. Department of Chemical Engineering. University of Illinois. (2002). <http://brahms.scs.uiuc.edu/>
- [19] Kano, M., Nagao, K., Hasebe, S., Hashimoto, I., Ohno, H., Strauss, R. and Bakshi, B. *Comparison of multivariate statistical process monitoring methods with applications to the Eastman challenge problem*. Computers and Chemical Engineering. 26. 161-174. (2002).
- [20] Mika, S., Rätsch, G., Weston, J., Schölkopf, B. and Müller, K. *Fisher discriminant analysis with kernels*. Neural Networks for Signal Processing IX, Y.-H. Hu et al., Eds. New York. IEEE, pp.41-48. (1999).
- [21] He, P., Wang, J. and Qin, J. *A new fault diagnosis method using fault directions in Fisher discriminant analysis*. AIChE Journal, Vol.51, No.2, 555-571, (2005).

- [22] Hernández, J., Ramírez, J. y Ramírez, C. *Introducción a la minería de datos*. PEARSON EDUCACIÓN, S.A., Madrid. España. (2004).
- [23] Johannesmeyer, M. *Abnormal situation analysis using pattern recognition techniques and historical data*. M.Sc. Thesis, University of California, Santa Barbara, CA. (1999).
- [24] Johannesmeyer, M., Singhal, A. and Seborg, D. *Pattern matching in historical data*. AIChE Journal. Vol. 48, No. 9, 2022-2038. (2002).
- [25] Kano, M., Tanaka, S., Hasebe, S. and Hashimoto, I. *Monitoring independent components for fault detection*. AIChE Journal. Vol. 49, No. 4, 969-976. (2003).
- [26] Chen, J. and Howell, J. *Towards distributed diagnosis of the Tennessee Eastman process benchmark*. Control Engineering Practice 10. pp. 971-987. (2002).
- [27] Chiang, L., Kotanchek, M. and Kordon, A. *Fault diagnosis based on Fisher discriminant analysis and support vector machines*. Computers and Chemical Engineering. 28. 1389-1401. (2004).
- [28] Guo, M., Xie, L., Wang, S. and Zhang, J. *Research on an integrated ICA-SVM based framework for fault diagnosis*. IEEE International Conference on Systems, Man and Cybernetics. (2003).
- [29] Smith, C. and Corripio, A. *Principles and Practice of Automatic Process Control. Second Edition*. John Wiley & Sons, Inc. USA. (1997).
- [30] Choudhury, M.A.A.S., Thornhill, N. and Shah, S.L. *Modelling valve stiction*. Control Engineering Practice, 13 (5). pp. 641-658. (2005)