

Diagnóstico de fallas en procesos técnicos: un enfoque basado en minería de datos y SVM

Fault diagnosis in technical process: a data mining and SVM setting

Ríos-Bolívar Addison¹; Hidrobo Francisco²; Guillén Pablo³

¹Departamento de Control, Facultad de Ingeniería, Universidad de los Andes

²SUMA, Facultad de Ciencias, Universidad de los Andes

³CESIMO, Facultad de Ingeniería, Universidad de los Andes

Mérida, - 5101 Venezuela

*ilich@ula.ve

Abstract

El mayor inconveniente para la detección y diagnóstico de fallas en procesos técnicos basados en redundancia analítica es el requerimiento de un modelo muy preciso del sistema. Por el contrario, en los métodos basados en el manejo de datos no se requiere de un modelo preciso, ya que se fundamentan en la manipulación de la información por medio de los datos medidos. Así, en este trabajo se presentan técnicas de representación que permiten cuantificar la cantidad de información contenida en los datos recolectados de los sistemas de Monitoreo, Diagnóstico y Detección (MDD). La representación obtenida es usada para la reconstrucción de los patrones de fallas, y seguidamente para su clasificación, empleando una máquina de aprendizaje de vectores de soporte (SVM), a objeto de lograr el diagnóstico de las fallas. Para verificación de los resultados se emplean datos generados por dos modelos, uno lineal, con un observador de estados y otro modelo no-lineal que representa el control de un levitador magnético. Se comparan los resultados usando una representación basada en índices estadísticos y otra basada en una aproximación polinómica; en ambos casos se logra la reconstrucción de patrones, que permite separar los distintos comportamientos dinámicos (fallas) con lo que se obtiene una clasificación entre las fallas (clases).

Palabras claves: Monitoreo, diagnostico, detección, localización, clasificación

Resumen

The main drawback for faults detection and diagnosis in technical processes based on analytical redundancy is that we need a very accurate model of the system. In contrast, when we use methods based on data mining, we do not need an accurate model because they are based on manipulation of the information by means of measured data. Thus, in this paper we represent techniques to quantify the amount of information contained in the data collected from systems of monitoring, diagnostic and detection (MDD). The representation obtained is used to reconstruct faults patterns, and then for classification using a Support Vector Machine (SVM), in order to make the faults diagnosis. For results verification, we used two models, a linear model with a state observer and a non-linear model that represents the control of a magnetic levitation. Results are compared using a representation based on statistical indexes and a representation based on a polynomial approximation; in both cases we achieved the reconstruction of patterns, which allow us separating different dynamic behaviors (faults) as wells as obtain a classification between faults (classes).

Key words: Monitoring, fault detection and isolation, classification, localization.

1 Introducción

La confiabilidad operacional, en procesos de producción, debe estar conformada por: la correcta operación de

los procesos, los sistemas de control asociados y la coordinación de los mismos.

En el contexto de operación confiable y segura, se desarrollan los sistemas que permiten el reconocimiento de eventos, los cuales deben orientar la toma de decisiones

cuando el desempeño del proceso productivo se ve afectado por la presencia de cualquier eventualidad adversa. Puesto que la confiabilidad está muy ligada al concepto de seguridad, es fundamental de dotar a los procesos industriales de mecanismos de seguridad, cuyos elementos básicos son los sistemas de Monitoreo, Diagnóstico y Detección (MDD) (Ríos-Bolívar A., 2001).

Los sistemas de MDD se fundamentan en su capacidad para responder ante situaciones inesperadas, de manera que su principal tarea es la Detección y Diagnóstico de Fallas (DDF). Un sistema de DDF utiliza las mediciones del proceso a objeto de producir unos residuales, a partir de los cuales, mediante funciones de evaluación y lógicas de decisión, se busca la identificabilidad y la separabilidad de las fallas, con perspectivas de pronóstico y de mantenimiento autónomo.

Para la generación de residuos se aplican principios de redundancia, bien sea física o mediante modelos, (Patton y col., 2000). El abordaje mediante la redundancia analítica se enfoca en el uso de modelos característicos de los procesos, generalmente, modelos matemáticos. Los modelos basados en conocimiento se fundamentan más en la información o datos disponibles, lo que conlleva a disponer de mecanismos caracterizados por la inteligencia artificial.

Por otro lado, se pueden generar residuos a través del manejo de datos, lo cual conlleva a la aplicación de métodos numéricos y otras técnicas de la inteligencia artificial como el aprendizaje de máquinas. Siguiendo este contexto, la taxonomía para la generación de residuos se puede ampliar al considerar los métodos basados en el manejo de datos, (ver Fig. 1). Entonces, lo principal para el DDF es el manejo de la información de los procesos a los fines de producir los residuos.

Dentro de estos métodos analíticos, la utilización de un observador de estados ha sido ampliamente estudiada. La idea es la generación de residuales con propiedades direccionales precisas, mediante una selección adecuada de la ganancia del observador, (Ríos-Bolívar 2003).

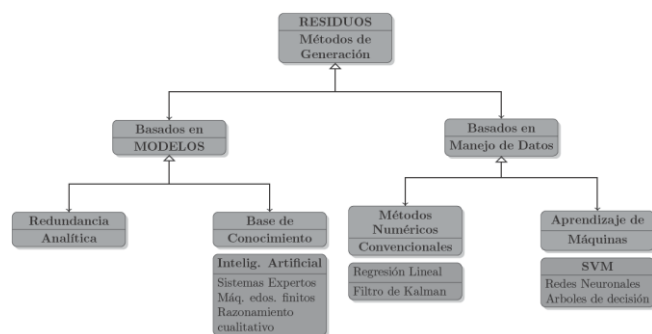


Fig. 1. Taxonomía con Algoritmos de Manejo de Datos.

Por otro lado, en muchos procesos industriales, desafortunadamente, los efectos de las fallas, las señales de ruido y las incertidumbres del modelo no se pueden separar

uno del otro. En estos casos se busca ampliar la capacidad o el desempeño de la detección al hacer que los residuos sean menos sensibles a las incertidumbres con respecto a una dirección de falla, en particular mediante el uso de técnicas de estimación óptima de estados, (Blanke y col., 2003). Fundamentalmente, estos son métodos de estimación robusta con el objetivo primordial de la supresión de la perturbación efectiva, manteniendo una sensibilidad adecuada con respecto a las fallas, (Ríos-Bolívar y Acuña, W., 2009). Dado que estas divergencias son proclives a la generación de residuos en ausencia de una falla verdadera, también se debe considerar, en el marco de las decisiones lógicas, un nivel de aseguramiento de la presencia de una falla, es decir, bajo la presencia de incertidumbres y de perturbaciones externas, a partir de los residuos debe existir un nivel de disparo (umbral) que garantice la presencia de una falla. Eso significa que a partir de los residuos, debemos disponer de un mecanismo o generador de función de decisión que nos permita distinguir, a partir de un umbral, cuando está presente un comportamiento anormal.

En base a las consideraciones anteriores, que debilitan la aplicación de métodos de redundancia analítica, se ha orientado a la utilización de métodos basados en datos. Algunos intentos iniciales orientados a la construcción de filtros inteligentes para la detección de fallas se pueden evaluar en (Ragot y col., 2000), (Ríos-Bolívar 2003). Fundamentalmente, se han utilizado técnicas de reconciliación de datos y redes neuronales.

Las Máquinas de Vectores de Soporte (SVM), son máquinas de base lineal y a solución única, fundadas en la teoría del aprendizaje estadístico. Estas máquinas aprenden la superficie de decisión de dos clases distintas de los puntos de entrada. Como un clasificador de una sola clase, la descripción dada por los datos de los vectores de soporte es capaz de formar una frontera de decisión alrededor del dominio de los datos de aprendizaje con muy poco o ningún conocimiento de los datos fuera de esta frontera. Los datos son mapeados por medio de un kernel Gaussiano, u otro tipo de kernel, a un espacio de características en un espacio dimensional más alto, donde se busca la máxima separación entre clases (Burges C. y Schölkopf B, 1997), (Burges 1997, 1998). En consecuencia, las SVM resultan en un excelente algoritmo para clasificación y reconocimiento de patrones, exhibiendo ciertas mejoras respecto a las redes neuronales (Okba y col., 2010). Así, las SVM representan un mecanismo para la generación de residuos orientados a la detección y diagnóstico de fallas (Jack y col., 2002), (Batur y col., 2002), (Mahadevan y col., 2009, Xu y col., 2010). En ese contexto, en (Ríos-Bolívar y col., 2013a) y (Ríos-Bolívar y col., 2013b) se presentan aplicaciones de diagnóstico de fallas a partir de minería de datos y SVM. Allí se indica una manera de utilizar estas técnicas para la detección de fallas mediante el manejo de datos del proceso. A pesar de estos resultados, aún persisten problemas abiertos relativos al diseño de filtros de DDF basados en SVM con propiedades de robustez frente a perturbaciones e incer-

tidumbres, la detección de fallas no reconocidas a priori, un algoritmo eficiente para la identificación de fallas en línea, entre otros aspectos. De igual manera, el diseño y construcción de filtros de DDF que combine, integral y eficientemente, las fortalezas de las técnicas basadas en redundancia analítica y los métodos que consideran las SVM.

2 Filtros de Detección y Diagnóstico de Fallas

En referencia a los métodos analíticos, los sistemas dinámicos se pueden describir por modelos que entran dentro de dos categorías: los modelos representativos y los modelos de diagnóstico. Los modelos representativos permiten describir el comportamiento dinámico de los sistemas en términos de una estructura estandarizada que, de manera satisfactoria, se aproxima al comportamiento entrada-salida o al comportamiento de estado de los sistemas, en ellos se encuentran los modelos heurísticos. Por el contrario, los modelos de diagnóstico permiten describir el comportamiento dinámico a través de una réplica de la estructura o arquitectura física de los sistemas. Allí, las unidades funcionales básicas, o de interés, se modelan en forma explícita: modelos de sensores, de actuadores, etc. A los fines de diseñar los filtros de DDF, los modelos de diagnóstico son los más útiles, mientras que para propósitos de control, lo son los modelos representativos.

2.1 Filtros de DDF basados en observadores

Dado que en esta técnica analítica de diseño de filtros de DDF es fundamental un modelo matemático, consideremos la dinámica de un sistema lineal, continuo e invariante en el tiempo, LTI,

$$\begin{aligned}\dot{x}(t) &= Ax(t) + Bu(t), \quad x(0) = x_0, \\ y(t) &= Cx(t).\end{aligned}\quad (1)$$

donde los estados $x \in \mathfrak{X} \subset \mathbb{R}^n$, los controles $u \in \mathfrak{U} \subset \mathbb{R}^m$, las salidas $y \in \mathfrak{Y} \subset \mathbb{R}^q$; y las matrices $A \in \mathbb{R}^{n \times n}$, $B \in \mathbb{R}^{n \times m}$, $C \in \mathbb{R}^{q \times n}$. Allí, el par (A, C) es observable.

En la descripción del sistema en (1) es fácil distinguir tres subsistemas: la estructura del proceso representada por A ; los actuadores que se describen por B ; y los sensores representados por C . En cualquiera de esos subsistemas pueden presentarse situaciones de comportamiento inadecuado con respecto a las características de desempeño establecidas en el diseño. De modo que se deben construir los esquemas de monitoreo que generen los residuales en función de las posibles fallas en cada una de las unidades funcionales.

El diseño de filtros de DDF se puede dividir en dos etapas: la primera fase es la generación de los residuales, (el problema de detección). La segunda etapa es la evaluación de los residuales a objeto de determinar el origen de las

fallas, (el problema de separación de las fallas). Los residuales se producen al comparar la salida estimada con la salida medida de la planta física. Así, para el sistema (1) existe una matriz de ganancia $D \in \mathbb{R}^{n \times q}$ de modo que el estimado $\hat{x}(t)$ del vector de estados $x(t)$ será la solución para la ecuación del observador de orden completo:

$$\begin{aligned}\dot{\hat{x}}(t) &= A\hat{x}(t) + Bu(t) + D(y(t) - C\hat{x}(t)), \\ \hat{y}(t) &= C\hat{x}(t).\end{aligned}\quad (2)$$

Allí, $\hat{y}(t)$ son las salidas estimadas y D la ganancia de realimentación del observador que se debe seleccionar adecuadamente. Esto es, D se selecciona de manera que $(A - DC)$ sea estable, y el error de estimación sea nulo, $e(t) = 0$. Puesto que para $t < t_0$, el proceso tiene un funcionamiento normal, es decir, no existen fallas, en ese momento el residual es aproximadamente igual cero. Cuando cualquier falla se hace presente, en $t \geq t_0$, el residual es distinto de cero y ello es propicio para la detección de las fallas.

Podemos observar, a partir de las diferentes representaciones de las fallas, que las mismas se pueden describir como entradas adicionales en la dinámica del proceso, además, de la susceptibilidad a la existencia de distintas fallas en un mismo subsistema, en adelante adoptaremos el siguiente modelo de diagnóstico para los sistemas sometidos a fallas:

$$\begin{aligned}\dot{x}(t) &= Ax(t) + Bu(t) + \sum_{i=1}^k L_i \nu_i(t), \quad x(0) = x_0 \\ y(t) &= Cx(t) + \sum_{i=1}^k M_i \nu_i(t),\end{aligned}\quad (3)$$

donde: L_i , M_i son las direcciones de fallas en los sensores y actuadores, respectivamente. Se supone que estas direcciones son conocidas. En lo sucesivo, se supone que las direcciones de fallas son linealmente independientes.

$\nu_i \in \mathfrak{V}_i$, corresponde al modo de la falla. Es una función autónoma, arbitraria y desconocida. Además, $\nu_i = 0$ si $t < t_0$; y $\nu_i \neq 0$ para $t \geq t_0$.

k es el número de fallas que funcional y estadísticamente son significativas.

Bajo esta representación, donde existen múltiples fallas posibles, con diferentes direcciones, el problema de separabilidad de las fallas se hace evidente, debemos construir un generador de residuales que nos permita distinguir a cual dirección del espacio de estado de (3) pertenece la falla que se ha hecho presente.

En el caso de filtros basados en observadores, debemos construir una ganancia del observador de modo que el vector de residuales, esto es, la salida del error de estimación, tenga características en una sola dirección asociada con al-

guna dirección de falla conocida, el resultado es la separabilidad de las fallas en el espacio de las salidas. Como la información importante está en la dirección de la falla, no se requiere del conocimiento del modo de la falla.

Consideremos un observador de estados como en (2), para el sistema (3), entonces, la dinámica del error de estimación estará regida por

$$\begin{aligned}\dot{e}(t) &= (A - DC)e(t) + \sum_{i=1}^k (L_i - DM_i) \nu_i(t) \\ \eta(t) &= Ce(t) + \sum_{i=1}^k M_i \nu_i(t).\end{aligned}\quad (4)$$

con $e(0) = x_0 - \hat{x}_0$

Si D se selecciona de modo que $(A - DC)$ sea estable, y si $\nu_i(t) \neq 0$, entonces $e(t) \neq 0$, por lo tanto se producen los residuales, puesto que $\eta(t) \neq 0$. Cualquier cambio en el proceso, por efecto de fallas, es notoriamente acentuado en la innovación de la salida del observador, de este modo se completa la fase de generación de residuales. Por otro lado, debido a las propiedades de los observadores, el desacoplamiento en las condiciones iniciales no tiene mayor efecto.

En este momento, si la única condición que se impone para la selección de la matriz de ganancia del observador, D , es que $(A - DC)$ sea estable, no se está en capacidad para establecer una clara distinción entre los efectos de las diferentes fallas. En principio, se pudiese pensar en el diseño de un conjunto de observadores, cada uno de los cuales se hace corresponder con una dirección de falla específica, lo que sugiere el diseño de D_i , $i = 1, \dots, k$, ganancias lo cual no es ni práctica, ni elegantemente factible. La idea es construir un único filtro de DDF. De lo anteriormente expuesto resaltan dos preguntas importantes: ¿Cuándo una falla es detectable?, el problema de detección, y ¿Cuándo las fallas son separables?, el problema de diagnóstico. Las condiciones y algunas soluciones para estos problemas pueden ser evaluadas en (Ríos-Bolívar A 2001) y (Ríos-Bolívar 2003).

3. Aprendizaje de Máquinas y SVM

La teoría de las SVM fue desarrollada por Vapnik basado en la idea de minimización del riesgo estructural (Burges C., 1998). Primero, la SVM mapea los puntos de entrada a un espacio de características de una dimensión mayor (ejemplo, si los puntos de entrada están en $\mathbb{R}^2$ entonces serán mapeados a $\mathbb{R}^3$ por la SVM), para luego encontrar el hiperplano que los separe y maximice el margen m entre las clases, (Burges C., 1998). En consecuencia, una SVM tiene la estructura de una red estática basada en kernels, para rea-

lizar clasificación lineal sobre vectores transformados a un espacio de dimensión superior, es decir, separa mediante un hiperplano en el espacio transformado.

Así, las operaciones de una SVM son:

- Transformar los datos a un espacio de dimensión más alta, a través de una función kernel, lo que implica que se reformula el problema de tal forma que los datos se mapean implícitamente en este espacio.
- Encontrar el hiperplano que maximiza el “margen” entre dos clases, mediante el cálculo eficiente del hiperplano óptimo.
- Si los datos no son linealmente separables, encuentra el hiperplano que maximiza el margen y minimiza una función del número de clasificaciones incorrectas (término de penalización de la función).

Cuando los conjuntos son linealmente separables, se debe seleccionar el hiperplano que maximiza el margen m , ver Fig. 2.

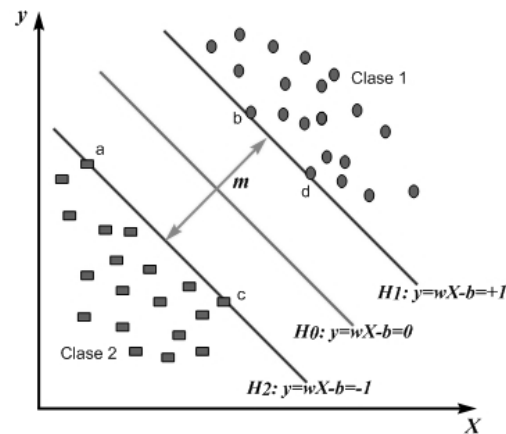


Fig. 2. Problema de separación lineal: Maximización del margen m .

Al maximizar el margen m , la separación de las clases es un problema de programación cuadrática, y puede ser resuelto por su problema dual introduciendo multiplicadores de Lagrange (Burges 1998). La SVM puede encontrar el hiperplano óptimo utilizando el producto escalar con funciones en el espacio de características que son los kernels. Los vectores de soporte son combinaciones de unos pocos puntos de entrada que permiten escribir la solución del hiperplano de manera más sencilla. Por ejemplo, considérese:

Un conjunto de N puntos de datos de entrenamiento: $\{(X_1, y_1), \dots, (X_N, y_N)\}$

Un hiperplano; $H_0 : y = wX - b = 0$

Donde: w es normal al hiperplano, $b / \|w\|$ es la distancia perpendicular al origen y $\|w\|$ es la norma euclídea de w .

Dos hiperplanos paralelos a H_0 :

$$\begin{aligned} H_1: y &= wX - b = +1 \\ H_2: y &= wX - b = -1 \end{aligned} \quad (5)$$

Con la condición de que no hay puntos de datos entre H_1 y H_2 .

Tal como se ilustra en la Fig. 2, si la distancia d_+ (d_-) es la distancia más corta desde la separación del hiperplano H_0 al punto más cercano positivo (negativo), donde el hiperplano H_1 (H_2) está ubicado, entonces la distancia entre los planos H_1 y H_2 es $d_+ + d_-$. Así, $d_+ = d_- = 1/\|w\|$ entonces el margen es igual a $2/\|w\|$. El problema es encontrar el hiperplano que dé el máximo margen. Los parámetros w y b son llamados vector peso y sesgo, respectivamente. La optimización del problema se presenta con la ecuación:

$$\min_{w,b} \frac{1}{2} w^T w \quad \text{restringido por} \quad y_i = wX - b \geq 1$$

La optimización del problema presentado en la ecuación anterior puede ser declarada como un problema convexo, cuadrático en (w, b) en un conjunto convexo. Usando la formulación del Lagrangiano, las limitaciones pueden ser reemplazadas por limitaciones de multiplicadores de Lagrange en sí mismos. Así, con un kernel adecuado, la SVM puede separar en el espacio característico los datos que en el espacio original de entrada no es separable (Burgess 1998).

En resumen, con el fin de separar un conjunto de datos se selecciona un conjunto de datos de entrenamiento (X, Y) , el problema de optimización es resuelto y los parámetros óptimos son calculados. Entonces, un vector de datos X dado del conjunto de datos iniciales es clasificado de acuerdo al valor de $(wX^* + b)$. La eficiencia de los vectores de soporte calculados es probada usando el conjunto de la data de prueba derivados del conjunto de datos iniciales.

4. Detección y Diagnóstico de Fallas con SVM

La técnica consiste en la generación de residuos, para el problema de detección de fallas, a partir de filtro basado en observador de estados. Luego, para el diagnóstico de las fallas, los residuos son procesados para reconocimiento y clasificación de patrones usando SVM.

En el problema de DDF se ha mostrado la potencialidad de las SVM para la generación de residuos a través del reconocimiento y clasificación de patrones (Jack y col., 2002, Batur y col., 2002, Mahadevan y col., 2009). La solución que se alcanza consiste en determinar patrones de condiciones de operación normal de los procesos a través del entrenamiento de SVM. En condiciones de operación anómala, es decir, bajo fallas, se determinan diferencias (residuos) entre el modelo de comportamiento obtenido por medio de un filtro de DDF y el funcionamiento real del proceso, los cuales son procesados por una SVM para la separación de las fallas (Ríos-Bolívar y col., 2012).

Así, cuando múltiples fallas son susceptibles de aparecer, las SVM permiten el reconocimiento y clasificación de

patrones, por lo que se obtiene el diagnóstico de las fallas.

De manera esquemática, la Fig. 3 muestra los diferentes elementos a construir para obtener el sistema de DDF. Los residuos son generados con filtros observadores de acuerdo a la naturaleza del proceso, principalmente para sistemas lineales. Esto tiene la ventaja de reconocer, inmediatamente, la presencia de fallas.

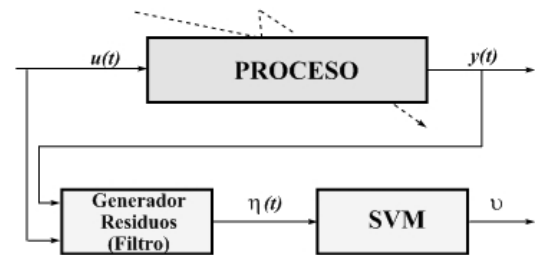


Fig. 3. Modelo para DDF basado en SVM

Para lograr el diagnóstico de las fallas mediante la reconstrucción de los patrones, se utiliza una SVM para procesar directamente los datos o residuos del filtro.

En ciertos casos, y principalmente para los sistemas no lineales, se torna complicado la generación de los residuos mediante observadores, siempre que las condiciones de detectabilidad no se satisfagan. Entonces es posible recurrir a las SVM para reconstruir ciertos patrones de fallas a partir de las entradas y salidas del proceso, ver Fig. 4.

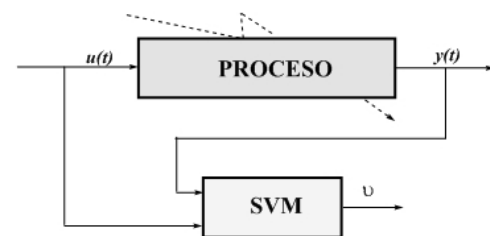


Fig. 4. Filtro de DDF basado en SVM

La reconstrucción de los patrones de fallas se realiza mediante la clasificación de patrones, donde la SVM utiliza directamente las entradas y salidas del proceso. No se generan residuos, por el contrario, se obtiene un estimado de los patrones de las fallas a partir de los cuales se obtiene el diagnóstico. Además, para evitar falsas alarmas o un diagnóstico incorrecto, se hace una minería de los datos a través de índices estadísticos y coeficientes polinomiales, los cuales se describen a continuación.

4.1 Representación de datos

Para la clasificación de fallas se usa una SVM que trabaja como una función que mapea los valores de entrada a un espacio compuesto por los estados de operación del sis-

tema. En este caso, dicha función recibe como entradas los valores medidos del sistema y produce como salida el estado del sistema (normal, en falla 1, en falla 2, etc.). Para el caso de sistemas dinámicos, las entradas a la SVM son un grupo de series temporales, una para cada valor medido; por la tanto se requiere una representación de estos datos que permita reconstruir el comportamiento de cada valor (señal).

Para tal representación se han estudiado dos esquemas:

Índices estadísticos. Se definen dos índices estadísticos para representar los datos en la ventana de tiempo:

Tamaño de la Curva. Esta característica es útil para conocer la estabilidad de los valores de una señal. Si en un intervalo, el valor de esta característica es bajo indica que la señal es estable, en caso contrario, la señal es inestable. La ecuación que define esta medida es:

$$L = \sum_{i=1}^{n-1} |x_{i+1} - x_i|$$

Donde: n es el tamaño de la ventana de tiempo y los x_i son los valores en esa ventana.

Umbral. La determinación del umbral γ se basa en el cálculo de la desviación de los datos para saber que tan dispersos están en una ventana de tamaño de n . El umbral se determina mediante:

$$\gamma = \frac{3}{n-1} \sqrt{\sum_{i=1}^n (x_i - \bar{X})^2}$$

Coefficientes polinomiales.

La idea es buscar ciertos valores que reproduzcan el comportamiento de la señal en una ventana de tiempo, obteniendo un conjunto de valores de menor tamaño que la ventana. Así, se define un polinomio de grado 3, con la forma: $ax^3 + bx^2 + cx + d$.

El procedimiento consiste en tomar los valores de la ventana y pasar un polinomio de interpolación de grado 3; esto genera como resultado los valores a , b , c y d para la ventana de tiempo. De esta manera, se obtienen 4 coeficientes, independientemente del tamaño de la ventana; estos coeficientes son usados como entrada a la SVM, para cada variable medida.

5. Ejemplos numéricos

5.1 Ejemplo de DDF para Sistemas Lineales

Considérese el modelo de diagnóstico lineal descrito por:

$$\begin{aligned} \dot{x} &= \begin{pmatrix} 0 & -1 & -2 \\ -1 & 0 & 1 \\ 2 & 1 & 0 \end{pmatrix} x + \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} u + \\ &\quad \begin{pmatrix} 0 \\ 1 \\ 0 \end{pmatrix} \nu_1 + \begin{pmatrix} 1 \\ 0 \\ 1 \end{pmatrix} \nu_2 \\ y &= \begin{pmatrix} 1 & 0 & 0 \\ 0 & 0 & 1 \end{pmatrix} x \end{aligned}$$

Para este sistema de diagnóstico, a pesar de que satisface las condiciones de detectabilidad y separabilidad para las fallas ν_1 y ν_2 , no es posible diseñar un filtro de DDF basado en observadores que asegure, simultáneamente, la estabilidad asintótica del error de estimación (detección) y la separación de las fallas, solo se obtiene la detección de las fallas (Szigeti col., 2000).

En ese sentido se ha diseñado un filtro con ganancia:

$$D = \begin{pmatrix} 0 & -2 \\ -11 & 11 \\ 0 & 10 \end{pmatrix}$$

de tal manera que la matriz $A - DC$ es estable. De esta manera se resuelve el problema de detección de fallas.

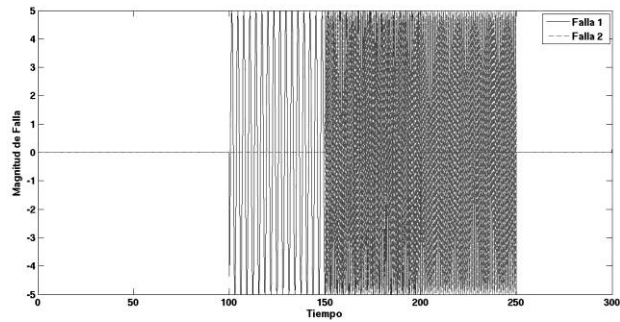
Como ha sido mencionado, sobre la base de las condiciones de detectabilidad y separabilidad de fallas, se pueden reconstruir los modos de fallas aplicando modelos basados en datos. Entonces, para la reconstrucción de los patrones de fallas se recurre a las SVMs.

A partir de un filtro de detección se generan los residuos, los cuales, en conjunto con las variables de control y la salida medida, se utilizan para el entrenamiento de una SVM. A continuación se realizan las fases de validación y verificación de manera de garantizar el desempeño del modelo de diagnóstico.

Para comprobar la efectividad de las SVM, se simula el sistema lineal por 300 segundos con cuatro estados posibles:

- 1: Normal. No se presentan fallas. ($t < 100$ y $t \geq 250$)
- 2: Falla1. Está presente la falla ν_1 . ($100 \leq t < 150$). Esta falla tiene un comportamiento sinusoidal con amplitud 5 y frecuencia 2.
- 3: Falla2. Está presente la falla ν_2 . ($200 \leq t < 250$). Esta falla tiene un comportamiento sinusoidal con amplitud 5 y frecuencia 15.
- 4: Múltiples Fallas. Las dos fallas presenten. ($150 \leq t < 200$)

La Fig. 5 muestra el comportamiento temporal de las dos fallas.

Fig. 5. Comportamiento temporal ν_1 y ν_2 .

Con los datos generados se entrena una SVM, usando un kernel gaussiano. Se usan como valores de entrada el control, las salidas del sistema y el error; este último calculado como la diferencia de la salida del sistema y la salida esperada obtenida a través del observador. Se toman una ventana de tiempo de tamaño 10, como la simulación se hace tomando el paso a 0,05 segundos; esto representa una ventana de 0,5 segundos. Con el objeto de reproducir la tendencia de las series temporales en cada ventana y suavizarlos, se usan los dos métodos descritos anteriormente, los índices estadísticos y los coeficientes polinomiales. En ambos casos se obtiene una clasificación exacta en la fase de entrenamiento (0% de error).

Caso 1: para comprobar los resultados de la SVM en un escenario diferente, se usaron los mismos parámetros para las fallas, pero se modificó el sistema para producir éstas en distintos instantes de tiempo:

- : Normal. No se presentan fallas. ($t < 100$ y $t \geq 250$)
- : Falla1. Está presente la falla ν_1 . ($200 \leq t < 250$)
- : Falla2. Está presente la falla ν_2 . ($100 \leq t < 150$)
- : Múltiples Fallas. Las dos fallas presentes. ($150 \leq t < 200$).

La Fig. 6 y la Fig. 7 muestran el resultado que se obtiene con los índices estadísticos y los coeficientes polinomiales. Como puede verse, usando índices estadísticos se logra un mejor seguimiento (predicción) del estado del sistema, y por tanto un mejor diagnóstico de las fallas con relación al uso de coeficientes polinomiales. Para este caso, con los índices estadísticos se obtienen predicciones correctas para el estado del sistema en más del 99 % del tiempo; mientras para los coeficientes polinomiales este valor está por debajo del 90 %.

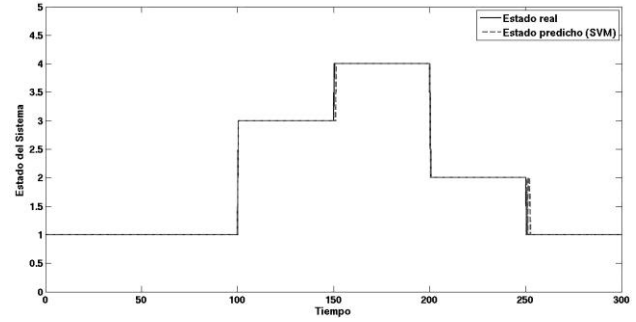


Fig. 6. Estado real y estado predicho con SVM e índices estadísticos

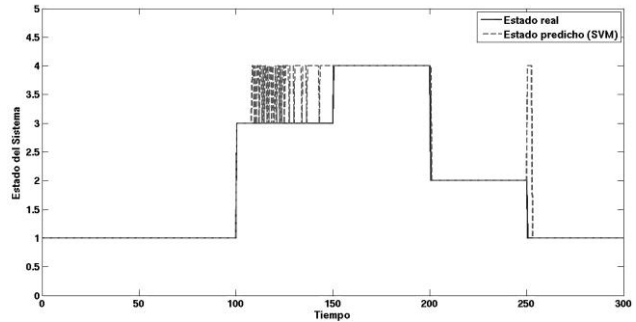


Fig. 7. Estado real y estado predicho con SVM y coeficientes polinomiales

Caso 2: para verificar la calidad del entrenamiento, se modificó el comportamiento de las fallas, pasando la frecuencia de la ν_1 de 2 a 2.5 y la amplitud de la ν_2 de 5 a 6.

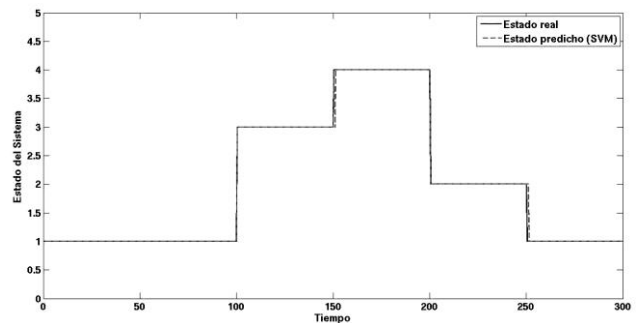


Fig. 8. Estado real y estado predicho con SVM e índices estadísticos.

Para el segundo caso, las Figuras 8 y 9 muestran el estado del sistema (real y predicho) con índices estadísticos y con coeficientes polinomiales. De nuevo, se observa que con índices estadísticos se logra una mejor clasificación de la falla (más de 99% de éxito en la predicción); mientras que los coeficientes polinomiales la predicción no es adecuada, obteniéndose un error superior al 30 %.

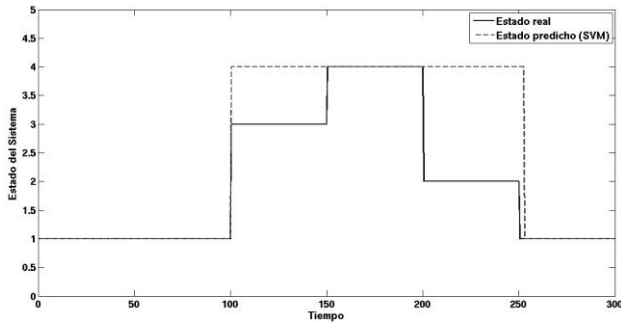


Fig. 9. Estado real y estado predicho con SVM y coeficientes polinomiales

5.2 Ejemplo para el Levitador Magnético

Un modelo del sistema de levitación magnética está dado por, (Lee S.H y col., 2000)

$$\begin{aligned}\dot{x}_1 &= x_2 \\ \dot{x}_2 &= -\frac{k}{2m}(x_3)^2 + g + \nu_1(t) \\ \dot{x}_3 &= \frac{x_2 x_3}{x_1} - \frac{R}{k} x_1 x_3 + \frac{1}{k} u(t) + \nu_2(t) \\ y_1 &= x_1, \quad y_2 = x_3\end{aligned}$$

donde: $x_1 = z$ (distancia de separación vertical), $x_2 = \dot{z}$ (velocidad vertical relativa), $x_3 = i$ (corriente en el magneto), g (fuerza de aceleración de la gravedad), m (masa del objeto suspendido), $u(t)$ (voltaje a través del magneto) y $k = \frac{\mu_0 N^2 A}{2}$ (factor fuerza). Los parámetros físicos son: $A = 0.04 \text{ m}^2$, $N = 660$, $\mu_0 = 4\pi \cdot 10^{-7}$, $R = 1\Omega$, $g = 9.8 \text{ m/s}^2$, $m = 0.300 \text{ Kg}$.

En este ejemplo no se diseña un filtro para la generación de residuos, debido a la naturaleza no lineal del proceso. Se aplica la recolección de los datos según la Fig. 4. Al igual que en el caso del modelo lineal, se simula 300 segundos con 4 estados:

- 1: Normal. No se presentan fallas. ($t < 100$ y $t \geq 250$)
- 2: Falla1. Está presente la falla ν_1 . ($100 \leq t < 150$). Esta falla tiene un comportamiento sinusoidal con amplitud 0.5 y frecuencia 1.
- 3: Falla2. Está presente la falla ν_2 . ($200 \leq t < 250$). Esta falla tiene un comportamiento sinusoidal con amplitud 15 y frecuencia 1, y se toma el valor absoluto de la misma.
- 4: Múltiple Falla. Las dos fallas presentes. ($150 \leq t < 200$)

En este caso, el comportamiento temporal de las fallas es mostrado en la Fig. 10.

Para la fase de entrenamiento, al usar los coeficientes polinomiales no se comete ningún error de entrenamiento; mientras que para los coeficientes estadísticos el error es menor al 2 %.

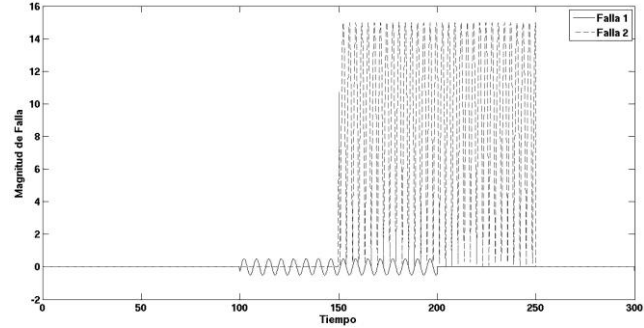


Fig. 10. Evolución temporal de las fallas

Se aplican las mismas técnicas del caso lineal. Para comprobar la efectividad del SVM en un escenario diferente, se usaron los mismos parámetros para las fallas, pero se modificó el sistema para producir éstas en distintos instantes de tiempo:

- 1: Normal. No se presentan fallas. ($t < 100$ y $t \geq 250$)
- 2: Falla1. Está presente la primera falla. ($200 \leq t < 250$)
- 3: Falla2. Está presente la segunda falla. ($100 \leq t < 150$)
- 4: Múltiple Falla. Las dos fallas presentes. ($150 \leq t < 200$)

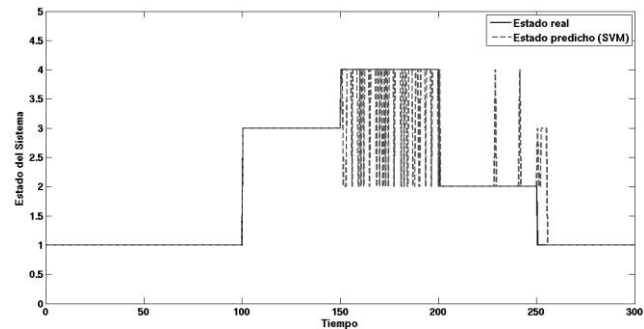


Fig. 11. Estado real y estado predicho con SVM e índices estadísticos

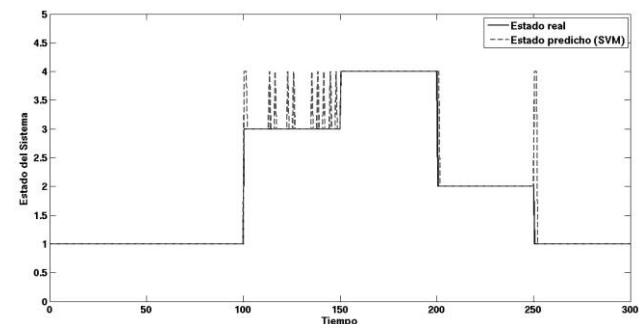


Fig. 12. Estado real y estado predicho con SVM y coeficientes polinomiales

A diferencia del ejemplo, para el levitador, a través de lo uso de los coeficientes polinomiales se logra una mejor clasificación en el caso 1; como puede apreciarse en las figuras 11 y 12. Con los coeficientes polinomiales se hace un mejor seguimiento del estado del sistema; se obtiene un error de predicción por debajo del 3 %; mientras que con los índices estadísticos este error es superior al 7 %.

Para una segunda evaluación, además de cambiar los tiempos en los que se producen las fallas, se modifican algunos valores de éstas. A la falla1 se le modificó la frecuencia de 1 a 1.1 y a la falla2 se le modificó la amplitud de 15 a 20.

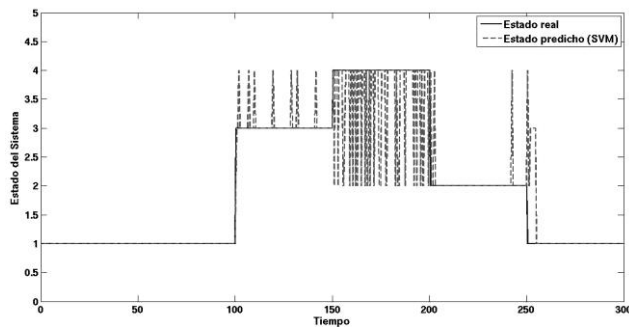


Fig. 13. Estado real y estado predicho con SVM e índices estadísticos

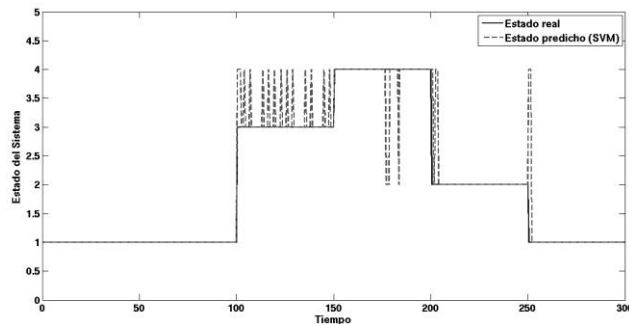


Fig. 14. Estado real y estado predicho con SVM y coeficientes polinomiales

Para este caso, la Fig. 13 y la Fig. 14 muestran la predicción que se logra para el estado del sistema usando índices estadísticos y coeficientes polinomiales, respectivamente. Como en el caso 1 de este ejemplo no lineal, con los coeficientes polinomiales la predicción es mejor con respecto a los índices estadísticos, obteniéndose un error del 5% para el primero y de más de 8 % para el segundo.

6. Conclusiones

Es posible implantar mecanismos de diagnóstico de fallas usando métodos basados en datos; para esto, se puede contar con esquemas de generación de residuos, siempre que sea posible construir observadores de estados adecuados, como se demostró para el ejemplo lineal. Estos métodos, en conjunción con técnicas de aprendizaje de máqui-

nas, como las SVM, proveen de herramientas para la eficaz clasificación de los estados de operación de un sistema. Como se observó, a través del ejemplo lineal, se puede reducir el volumen de datos a manejar por las SVM por medio del uso de índices estadísticos o coeficientes polinomiales; siendo los primeros más eficientes para el caso lineal. Además, se ha comprobado la efectividad de la combinación de técnicas basadas en redundancia analítica (filtro-observador), y técnicas basadas en datos (SVM), lo cual permite desarrollar sistemas de detección y diagnóstico de fallas factibles de ser implantados en procesos monitoreados y supervisados a tiempo real.

Por otro lado, estas técnicas aún pueden ser aplicadas cuando no sea factible construir un observador para el sistema, como se mostró en ejemplo del levitador magnético (sistema no lineal); obteniéndose resultados con alto índice de confiabilidad, por arriba del 95 %. Esta técnica puede ser mejorada por un lado, mediante la evaluación de otros kernels, y por otro, a través del reconocimiento de patrones por la vía de regresión, a objeto de reconstruir las señales de fallas por aproximación inversa.

7 Agradecimientos

Este trabajo ha sido financiado por el CDCHTA de la Universidad de Los Andes, a través del proyecto No. I-1268-11-02-A, por lo que gratamente se reconoce este apoyo.

Referencias

- Batur C, Ling Z, Chien-Chung C, 2002, Support vector machines for fault detection. Proc. 41st IEEE Conf. on Decision and Control, vol. 2, pp. 1355 – 1356.
- Blanke M, Kinnaert M, Lunze J, Staroswiecki M, 2003, Diagnosis and fault tolerant control. Springer-Verlag, Berlin.
- Burges C, Schölkopf B, 1997, Improving the accuracy and speed of support vector machines. Adv. in Neural Information Proc. Systems, vol. 9, pp 375-381. MIT Press.
- Burges C, 1998, A tutorial on support vector machines for pattern recognition. Data Mining and Knowledge Discovery, vol. 2:121-167.
- Jack LB, Nandi AK, 2002, Fault detection using support vector machines and artificial neural networks, augmented by genetic algorithms. Mechanical Systems and Signal Processing, 16(2-3), pp. 373-390.
- Mahadevan S, Shah SL, 2009, Fault detection and diagnosis in process data using one-class support vector machines, Journal of Process Control, 19(10), pp. 1627-1639.
- Lee SH, Sung HK, Lim JT., Bien Z, 2000, Self-Tuning Control of Electromagnetic Levitation Systems. Control Engineering Practice, Vol. 8, pp. 749-756.
- Okba T, Ilyes T, Tarek G, Hassani M, 2010, Supervised learning with kernel methods, Proc. 10th WSEAS Int. Conf. on Wavelet analysis and multirate systems, WAMUS'10, pp 73-77, Wisconsin, USA.

Patton RJ, Frank PM, Clark RN, 2000. Issues of Fault Diagnosis for Dynamical Systems, Springer-Verlag.

Ragot J, Kratz F, Maquin D, 2000, Finite memory observer for input-output estimation: Application to data reconciliation and diagnosis. 4th IFAC Safeprocess Conference, pp 587-592, Budapest - Hungary.

Ríos-Bolívar A, 2001, Sur la Synthèse de Filtres de Détection de Défaillances, PhD thesis, Université Paul Sabatier, Toulouse, France.

Ríos-Bolívar A, 2003, Metodologías para la Construcción de Sistemas de Detección y Diagnóstico de Fallas. PhD thesis, Universidad de Los Andes, Mérida - Venezuela, 2003.

Ríos-Bolívar A, Acuña W, 2009, Robust fault detection in uncertain polytopic linear systems. Revista Técnica de Ingeniería. Vol. 32, No. 2, pp. 1-10, LUZ. Maracaibo, Venezuela.

Ríos-Bolívar A, Hidrobo F, Guillén P, (2012), Diagnóstico de fallas en procesos dinámicos: Un enfoque basado en SVM. V CIBELEC, Mérida.

Ríos-Bolívar A, Guillén P, Hidrobo F, Rivas F, 2013a, Data Mining for Fault Diagnosis in Dynamic Processes: An approach based on SVM, Proc. 15th Int. Conf. on Math. and Computational Methods in Sci. and Eng. Kuala Lumpur, Malaysia, pp. 98-103.

Ríos-Bolívar A, Hidrobo F, Guillén P, Rivas F, 2013b, Fault Diagnosis in Dynamic Processes: A Data Mining and SVM Application, WSEAS Int. J. of Systems Engineering, Applications and Development, vol. 7, No. 4, pp. 191-199.

Szigeti F, Ríos-Bolívar A, Tarantino R, 2000, Fault detection and isolation filter design by inversion: The case of linear systems. 4th IFAC SAFEPROCESS Conference, Budapest, Hungary, pp 379-384.

Xu X, Wang S, Li F, Wu Z, Sun W, 2010, Research on fault diagnosis based on wavelet packet multi-class classification SVM, Int. Conf. on Manufacturing Automation, pp. 164-167.

Recibido: 18 de febrero de 2013

Revisado: 12 de junio de 2014

Ríos-Bolívar Addison: Profesor Titular en el Departamento de Sistemas de Control de la Facultad de Ingeniería, Universidad de Los Andes. Ha publicado más de 300 artículos científicos en Revistas, Libros y Actas de conferencias. El es autor del libro: "Control de Sistemas Lineales: Realimentando la Salida". También, es co-autor de los libros intitulados: "Sistemas MultiAgentes y sus Aplicaciones en Automatización Industrial", e "Implementando Técnicas de Control Acotado: Un Enfoque Basado en Tolerancia a Fallos".

Hidrobo, Francisco: Profesor Titular de la Facultad de Ciencias de la Universidad de Los Andes. Clasificado en el Programa de Estímulo al Investigador (PEI), tanto del Observatorio Nacional de Ciencia, Tecnología e Innovación como el PEI de la Universidad de Los Andes. Trabaja en las áreas de Automatización, Computación Inteligente y Computación de Alto Rendimiento. Correo electrónico: hidrobo@ula.ve.

Guillén, Pablo: Profesor Asociado en el Centro de Simulación y Modelado (CESIMO), Universidad de Los Andes, y Research Assistant Professor en el Department of Computer Science, University of Houston, Houston, TX, USA. El Dr. Guillén ha realizado una investigación Postdoctoral en Ciencias Computacionales, en el marco del Program in Computational Science, University of Texas, El Paso, El Paso, TX, USA, (2013). Correo electrónico: pguillen@ula.ve.