

Agrupaciones de modelos locales con descripción externa. Aplicación a una planta de frío solar

Manuel R. Arahál* Amparo Núñez-Reyes* Ignacio Alvarado Aldea*
Francisco Rodríguez**

* Dpto. Ingeniería de Sistemas y Automática, Universidad de Sevilla, Camino de los Descubrimientos, s/n, 41092, España, (dt: arahal@esi.us.es)

** Dpto. de Lenguajes y Computación, Universidad de Almería, La Cañada de San Urbano, 04120, España.

Resumen: En este artículo se presenta y analiza un método de creación de modelos formados por agrupaciones de submodelos locales lineales. La principal novedad es que la ponderación empleada con los submodelos no es la habitual, basada en el estado o parte del mismo sino que se toman exclusivamente señales instantáneas de entrada y salida del sistema. Esta elección simplifica algunos aspectos de la creación del modelo, manteniendo intacta la capacidad de representación, siendo ésta comparada con otras técnicas. La simplificación aludida es importante pues acerca el método a la práctica industrial del control de procesos. La técnica de identificación resultante se ilustra mediante dos casos prácticos: un sistema simulado propuesto por Narendra y el sistema de captación de una planta real de producción de frío a partir de energía solar. En ambos casos se muestran los errores de generalización para la predicción a un paso y para la simulación usando gran cantidad de situaciones. Los resultados indican que es factible el uso del método propuesto como técnica simplificada aplicable en la industria. Copyright © 2009 CEA.

Palabras Clave: agrupaciones de modelos locales, identificación, planta solar de producción de frío

1. INTRODUCCIÓN

Las agrupaciones de modelos locales lineales (en adelante AMLL) aparecen por primera vez en (Johansen and Foss, 1993) como un medio para describir dinámicas no lineales mediante el uso de un conjunto de submodelos lineales. La idea del método es interpolar entre submodelos en función del punto de operación del sistema. La salida proporcionada por el modelo es una suma ponderada de las salidas de cada submodelo. Los pesos de ponderación no son fijos sino que varían en función del punto de funcionamiento. La función que proporciona el valor de los pesos en cada instante es estática, pero no lineal, por lo que el modelo resultante tampoco lo es. Los submodelos lineales representan el comportamiento del sistema con cierta bondad en torno a un punto de funcionamiento y por ello se les llama modelos locales lineales (MLL). El modelo que resulta de la ponderación de los MLL es global, queriendo con esto decir que tiene validez en todos los puntos de operación.

Las AMLL han sido aplicadas a procesos (Foss *et al.*, 1995; Johansen *et al.*, 2000; Hunt *et al.*, 2000) utilizándose en la mayoría de los casos la ponderación de submodelos en función del estado o de algunas componentes de éste (Johansen and Foss, 1995; Murray and Johansen, 1997) (el lector interesado puede acceder a una biblioteca de funciones para MATLAB que ayudan en el diseño de este tipo de modelos (Johansen and Foss, 1998)). El uso de AMLL mediante descripciones externas no está tan extendido existiendo algún trabajo de corte teórico como (Bontempi and Birattari, 2005). En este artículo se mostrará que el uso de la descripción externa conforma un entorno más cercano a la idea de identificación automática, permitiendo la creación de un algoritmo para ser usado en la

práctica industrial. Otra ventaja del método propuesto es que sistematiza la selección de regiones de operación. La técnica de identificación resultante se ilustra mediante dos ejemplos, uno de ellos es un sistema propuesto por Narendra (Narendra and Parthasarathy, 1990), monovariable y simulado, el otro es el sistema de captación de una planta real de producción de frío a partir de energía solar con una salida y varias entradas. En ambos casos la dinámica observada depende en gran medida del punto de funcionamiento por lo que demuestran la posibilidad de usar el método como técnica simplificada aplicable en la industria de procesos.

Existen evidentes conexiones entre las AMLL y otros medios de identificación como los conjuntos borrosos. Esta cuestión se trata brevemente en el siguiente resumen del estado del arte.

1.1 Motivación y estado del arte

Los modelos lineales son ampliamente usados en la industria debido, entre otras razones, a la facilidad de obtención a partir de datos experimentales. Cuando se controla un proceso en torno a un punto de funcionamiento la aproximación ofrecida por el modelo lineal puede ser suficiente para los propósitos prácticos de control. Esta elección conlleva un menor tiempo de puesta en marcha del lazo de control, proporcionando además resultados aceptables en muchos casos.

Hay situaciones en las que las desviaciones del modelo lineal con respecto a la dinámica observada no son desdeñables. En tales casos puede ser conveniente una mayor inversión en modelado, pasando del simple modelo lineal a otro tipo de modelo más complejo en la esperanza de poder diseñar mejores controladores. Una opción atractiva en estos casos es la de usar un

controlador basado en modelo como los controladores predictivos (CPBM) (Camacho and Bordons, 2004), en particular en entornos donde se pretende llevar a cabo la automatización total (de Prada, 2004), (Rivera, 2007), (Sarabia *et al.*, 2006), (Rodríguez *et al.*, 2007), (Ghraizi *et al.*, 2007) pues los beneficios obtenibles mediante optimización justifican el gasto adicional invertido en modelado.

En el camino que lleva de los modelos lineales a los que no lo son hay un primer paso que consiste en añadir no linealidades estáticas como saturaciones, zonas muertas etc. Los modelos resultantes son llamados modelos con no linealidades concentradas y tienen la ventaja de aunar la simplicidad de los modelos lineales con la validez global de los modelos obtenidos a partir de leyes fundamentales. Su campo de aplicación es muy amplio, requiriendo un pequeño esfuerzo adicional con respecto al modelado lineal. El análisis de este tipo de modelos ha dado lugar a numerosos estudios (Alamo *et al.*, 2006) que proporcionan fundamentos teóricos sólidos para su aplicación. La principal desventaja es que son modelos poco flexibles en el sentido de que no son capaces de acomodar distintas dinámicas como tiempos característicos de respuesta, tiempos muertos, etc. en función del punto de operación.

El siguiente paso en complejidad son las AMLL. Su funcionamiento consiste en combinar las salidas de varios submodelos lineales locales. Estas agrupaciones de submodelos dan lugar a modelos globales, es decir, con validez en un espacio de trabajo amplio, no sólo en torno a un punto de trabajo particular. Además el modelo resultante es no lineal pues la combinación de los submodelos lineales se basa en funciones de ponderación que no son lineales. Los modelos que integran la agrupación son obtenidos mediante identificación en torno a puntos de trabajo; es decir, cada modelo es local pues constituye una aproximación adecuada en torno a un punto de trabajo particular. De este modo las técnicas de selección e identificación para modelos lineales pueden aplicarse por lo que el incremento de complejidad práctico no es elevado, de hecho se incrementa linealmente con el número de submodelos. Además, el orden de los submodelos lineales afecta de una manera suave a los requisitos prácticos de la identificación (excitación y número de datos) y no se produce la temida explosión dimensional (Bellman, 1957) que afecta a las series de funciones como por ejemplo las redes de neuronas de una capa oculta (Moody and Darken, 1989). Finalmente es de señalar que el modelo global resultante admite hasta cierto punto un análisis frecuencial. Análisis que puede obtenerse de manera más simple que en el caso de modelos con no linealidades concentradas y muchísimo más que en el caso de redes de neuronas.

La formulación presentada en este artículo consiste en agrupaciones de submodelos que usan la descripción externa del sistema y que usan exclusivamente valores instantáneos de entrada y de salida. En el apartado siguiente, se presentan las ideas de los MLL que sirven de base para diseñar un algoritmo de creación de AMLL mostrado en el punto tercero. Este algoritmo o método sistemático requiere únicamente datos medidos experimentalmente del sistema aunque permite igualmente la incorporación de conocimiento previo. Se mostrará además que las capacidades representativas de las AMLL son superiores a las de los modelos con no linealidades concentradas y similares a las de las redes de neuronas artificiales y los modelos borrosos. En el punto cuarto se presentan resultados prácticos obtenidos al aplicar el algoritmo a dos plantas.

2. AMLL EN EL ESPACIO DE ENTRADA-SALIDA

Los modelos locales usados habitualmente consisten en descripciones internas. Este tipo de descripción no siempre es de fácil utilización debido a la posible existencia de componentes del estado no medibles. En este trabajo se utilizan exclusivamente las señales de entrada y de salida, las cuales se suponen accesibles. Se persigue de este modo poner a disposición de la práctica industrial un algoritmo capaz de proporcionar modelos de sistemas con posible comportamiento no lineal a partir de datos obtenidos de ensayos experimentales.

A fin de presentar el método se considera un sistema monovariante que exhibe un comportamiento alejado de la linealidad cuya entrada es una señal escalar u_Σ y la salida es otra señal escalar y_Σ . No existe problema alguno para extender los resultados al caso de múltiples entradas y salidas aparte del aumento de complejidad de la notación. Para cada punto de funcionamiento dado por el par $z = (u_\Sigma^f, y_\Sigma^f)$ se va a considerar un modelo de entrada salida en las variables $y = y_\Sigma - y_\Sigma^f$ y $u = u_\Sigma - u_\Sigma^f$ que representan variaciones en torno al punto de funcionamiento. Se está suponiendo aquí que para una entrada dada u_Σ^f existe un único valor de régimen permanente y_Σ^f tal que $|y_\Sigma^f| < \infty$, los casos en que esto no se cumple son los comportamiento integradores y con histéresis y se tratarán en el punto tercero.

Como modelo correspondiente al punto de funcionamiento se propone una expresión de tiempo discreto en la forma:

$$y(k) = m^t(k-1)\theta \quad (1)$$

que contiene un vector de medidas $m(k-1) = (y(k-1), \dots, y(k-na), u(k-d-1), \dots, u(k-d-nb))^t$, donde d representa el retardo puro en ese punto de funcionamiento y $\theta = (a_1, \dots, a_{na}, b_1, \dots, b_{nb})^t$ un vector de parámetros que contiene los coeficientes $a_i \in \mathbb{R} \ i \in \{1, \dots, na\}$ y $b_j \in \mathbb{R}, j \in \{1, \dots, nb\}$ del modelo de órdenes $na \in \mathbb{N}$ y $nb \in \mathbb{N}$.

Si se enumeran los puntos de funcionamiento mediante p variando en $[1, np]$ se tiene que cada modelo local corresponde a un valor del índice p . De este modo la ecuación (1) puede particularizarse para cada punto de funcionamiento. A fin de marcar la diferencia entre la salida del sistema y la salida de cada modelo se denotará a éstas mediante $\hat{y}_p$. La predicción a un paso dada por cada submodelo queda descrita mediante

$$\hat{y}_p(k) = m_p^t(k-1)\theta_p, \ p = 1, \dots, np \quad (2)$$

Obsérvese que se contempla la posibilidad de que cada submodelo local tenga retardos d_p y órdenes na_p y nb_p diferentes a los otros submodelos. El caso más simple consiste en que $na_p = nb_p = 1 \forall p$. Este caso corresponde a usar modelos locales lineales de primer orden más tiempo muerto. Como es sabido este tipo de modelos son habitualmente usados en la práctica industrial entre otros motivos por la sencillez de identificación usando la curva de reacción (Camacho and Bordons, 2004). Para este caso simple se tiene además que los distintos vectores m_p para $p = 1, \dots, np$ colapsan en el mismo vector $m(k-1) = (y(k-1), u(k-1-d))^t$.

La ecuación (2) puede usarse para realizar la predicción local a un paso y de este modo obtener $\hat{y}_{\Sigma,p}(k) = \hat{y}_p(k) + y_{\Sigma,p}^f$. Por otra parte, la predicción a un paso para el modelo global se

denota mediante $\hat{y}_\Sigma$ y se obtiene como media ponderada de las predicciones locales, de manera que

$$\hat{y}_\Sigma(k) = \sum_{p=1}^{np} \omega_p \hat{y}_{\Sigma_p}(k) \quad (3)$$

Los pesos de ponderación ω_p se calculan a partir del punto de operación usando unas funciones de pertenencia (nótese la analogía con los conjuntos borrosos). Las funciones de pertenencia $\rho(z)$ se evalúan en el espacio (Z) de los puntos de funcionamiento $z = (u_\Sigma^f, y_\Sigma^f)$ según la expresión:

$$\rho(z) = \prod_{j=1}^{nz} \sigma_j(z) \quad (4)$$

siendo $\sigma(z)$ una función meseta que vale uno si z está en el interior de un conjunto $\mathcal{R}$ que contiene al punto de funcionamiento z^f y que cae suavemente al valor cero a medida que z se aleja de dicha región. En notación matemática se tiene que

$$\sigma_j(z) = \begin{cases} 1, & \text{si } \underline{z} < z_j < \bar{z} \\ \psi(z_j - \underline{z}), & \text{si } z_j < \underline{z} \\ \psi(\bar{z} - z_j), & \text{si } z_j > \bar{z} \end{cases} \quad (5)$$

siendo $\underline{z} = z^f - \Delta_j$, $\bar{z} = z^f + \Delta_j$. El tamaño de la meseta en la cual la función σ vale uno viene dado por $2\Delta_j$. La función $\psi(x) = e^{-\beta x^2}$ cae suavemente hacia cero con una pendiente que depende del parámetro β .

A partir de las funciones de pertenencia para cada punto de operación se obtienen los pesos de ponderación ω_p simplemente normalizando las funciones de pertenencia mediante

$$\omega_p(z) = \frac{\rho_p(z)}{\sum_{h=1}^{np} \rho_h(z)} \quad (6)$$

de este modo se cumple que $\omega_p(z) < 1$ para todo z y para todo $p = 1, \dots, np$, además la suma de las funciones de ponderación vale $\sum_{p=1}^{np} \omega_p(z) = 1$ para todo z .

Una vez obtenidos los parámetros de los submodelos mediante identificación en torno a cada punto de operación considerado, se puede utilizar las ecuaciones (1)-(6) para obtener la predicción a un paso $\hat{y}_\Sigma(k)$. Obsérvese que basta con conocer la entrada y la salida de la planta en algunos instantes anteriores dependiendo de los órdenes na_p , nb_p y los retardos puros d_p . Debido a la elección de la función de ponderación la salida de cada submodelo p tiene un peso mayor sobre una región del espacio $\mathcal{R}_p \subset \mathcal{Z}$ próxima al punto de operación z_p^f .

La predicción multipaso con horizonte H implica obtener las predicciones del modelo global $\hat{y}_\Sigma(k+j)$ con $j = 1$ hasta $j = H$ usando las predicciones anteriores en lugar de los valores futuros, y por tanto desconocidos, de y .

3. ALGORITMO DE CREACIÓN DE AMLL

A continuación se presenta el algoritmo propuesto para la creación de AMLL. La versión que se muestra contempla el caso más general posible que es el que requiere menor intervención por parte del usuario. Hay que tener presente que algunos de los pasos del algoritmo pueden realizarse de forma diferente a la indicada si se dispone de información previa. Surgen de este modo atajos que pueden ahorrar tiempo al usuario y que permiten además la incorporación de conocimiento a priori.

3.1 Algoritmo básico

Se supone que se dispone de un conjunto de medidas realizadas sobre el sistema que se quiere modelar. Sean estas medidas los puntos $(u_\Sigma(t), y_\Sigma(t))$ con $t = 1, \dots, nm$. Aunque resulte obvio conviene recordar que es preciso asegurarse de que el número de medidas nm es suficientemente alto comparado con el número de parámetros que se van a identificar. En un apartado posterior se discutirá la influencia de las señales de entrada aplicadas al sistema. En particular interesa realizar un tratamiento de las mismas para eliminar falsas medidas, para acondicionar la media y varianza de las distintas señales, etc. Una vez satisfechos estos requisitos se puede utilizar el algoritmo que consta de los siguientes pasos:

- Definir el espacio de los puntos de funcionamiento. En el caso simplificado que se está presentando $z = (y, u)$. En el caso de sistemas multivariable con ne entradas y ns salidas es preciso construir ns modelos de una sola salida teniendo cada uno de ellos un espacio $z_j(k) = (y_j, u_1, \dots, u_{ne})$, para $j = 1, \dots, ns$.
- Decidir el número de puntos de operación np .
- Fijar los valores correspondientes a los puntos de operación y sus regiones asociadas. De este modo se obtienen $z_p^f = (u_p^f, y_p^f)$ para $p = 1, \dots, np$. Es interesante hacer notar que los puntos z_p^f deben caer sobre la característica estática $y_e = f(u_e)$ del sistema.
- Construir las funciones de pertenencia. Estas funciones tendrán la forma general dada por la ecuación (4). Los detalles particulares dependen de la forma en que se dividan las regiones del espacio $\mathcal{Z}$. Por ahora se va a indicar un método simple dejando para más adelante las posibles variaciones. La división en regiones se realiza indicando para cada una un centro y la anchura de la región asociada (parámetros Δ de las funciones σ). Conviene definir las siguientes variables auxiliares:
 - Anchura de las regiones. A partir de la localización de los centros se pueden obtener los incrementos posibles $\Delta_{u,p}$ y $\Delta_{y,p}$ aceptables para cada región asociada a cada punto de operación. Una vez hecho esto la región puede representarse mediante la caja $[u_p^f - \Delta_{u,p}, u_p^f + \Delta_{u,p}] \times [y_p^f - \Delta_{y,p}, y_p^f + \Delta_{y,p}]$
 - Parámetro β . Este parámetro permite que la función de pertenencia decrezca con mayor o menor rapidez a medida que el punto de trabajo se aleja del punto considerado para derivar el submodelo local. Es conveniente que exista cierto solape entre funciones de pertenencia de regiones adyacentes pues de este modo se consigue una interpolación suave entre modelos.
- Fijar los órdenes de los submodelos locales.
- Calcular los retardos para cada submodelo. La forma más adecuada consiste en realizar experimentos especiales destinados a obtener estos valores como por ejemplo ensayos en escalón. En algunos casos los retrasos debidos al transporte pueden ser estimados mediante consideraciones físicas calculando el tiempo de residencia de fluidos. En último lugar, esta tarea puede realizarse a partir de los datos usando técnicas de correlación. Sea cual sea el método empleado el resultado de este paso es el conjunto de retardos para cada modelo local d_p .
- Identificar los submodelos locales obteniendo los vectores de parámetros θ_p correspondientes a cada submodelo.

Mediante este algoritmo quedan fijados todos los parámetros de la AMLL. Por otra parte el algoritmo necesita ciertos valores que el usuario debe proporcionar. Estos valores son llamados hiperparámetros o parámetros del algoritmo como np y β . El resto de valores pueden obtenerse a partir de los datos aunque también existe la posibilidad de que el usuario desee fijar alguno de ellos como se comenta a continuación.

3.2 Variantes del algoritmo

Ya se indicó anteriormente que el algoritmo admite variaciones en algunos de sus pasos incluyendo entre éstas la posibilidad de usar atajos. Se van a comentar en este punto algunas de las más interesantes.

- Número de puntos de funcionamiento. En lugar de fijar np es posible determinar el error máximo admisible para el modelo global y a partir de ahí obtener np mediante aproximaciones sucesivas. Esto es posible en cualquier dispositivo de cálculo numérico actual. Es conveniente tener en cuenta que dadas unas medidas y un valor de np el error típico obtenido en la predicción con un modelo formado por agrupación de MLL es $\epsilon(np)$. La función $\epsilon(np)$ decrece siempre a medida que np aumenta. Por este motivo es preciso fijar un límite al número de modelos locales. Alternativamente es posible utilizar un conjunto de medidas no usadas por el algoritmo y contrastar el error obtenido por el modelo con este nuevo conjunto con el previsto por $\epsilon(np)$. Este procedimiento habitual en el entrenamiento de redes de neuronas es una forma simple de validación (Reed and Marks, 1998).
- Selección de regiones. La forma habitualmente usada en identificación neuronal y borrosa consiste en la partición del espacio en regiones. Esta forma de proceder conlleva un número excesivamente alto de regiones pues algunas de las particiones nunca son visitadas en la evolución de la planta.

Con respecto a este problema ha de tenerse en cuenta que la formulación externa conlleva una disminución de la dimensión del espacio a dividir, lo cual es una ventaja considerable.

El algoritmo presentado es la más simple de varias opciones posibles. Consiste en dividir el intervalo de variación de la señal de entrada en segmentos similares, calcular el valor de régimen permanente de la salida que le corresponde a cada extremo del segmento y considerar las regiones resultantes. Cabe indicar que, debido a que la ganancia estática del sistema no es necesariamente constante en el espacio de trabajo, el procedimiento indicado puede dar lugar a regiones muy disimilares. Esto da lugar a la primera variante que pretende dividir más finamente las zonas en las que la ganancia estática varía más acusadamente. Una forma de solucionar el problema consiste en dividir directamente la característica estática en np regiones similares del espacio $\mathcal{Z}$, asignando los centros de tales regiones a los valores z_p^f .

Esta variante reduce además en cierta medida el problema de considerar regiones en las que nunca hay muestras. En efecto, al disponer las regiones sobre la característica estática se está favoreciendo las regiones donde hay más densidad de muestras. El problema sin embargo no desaparece del todo y queda abierta la puerta a nuevas mejoras. En particular, las técnicas de agrupamiento permiten seleccionar un número predeterminado np de pun-

tos del espacio $\mathcal{Z}$ tales que representen la distribución observada en las medidas (Principe *et al.*, 1998). Estos puntos seleccionados pueden utilizarse como puntos de funcionamiento z_p^f , la división en regiones se lleva a cabo posteriormente por crecimiento de regiones de Voronoi (Fritzke, 1994) tomando como centros los puntos de funcionamiento.

La figura 1 ilustra los tres procedimientos considerados. Se ha utilizado como ejemplo un hipotético sistema dinámico de una entrada y una salida. En cada gráfica el eje horizontal corresponde a la entrada (e) y el eje vertical a la salida (s). Se ha trazado con línea discontinua la característica estática obtenida de observaciones experimentales. Para cada método de selección de regiones se muestra la frontera de cada región. Dicha frontera se ha calculado de forma que la zona interior produzca un valor mayor o igual al 90% para la función de pertenencia. De arriba a abajo se tienen los resultados obtenidos por división homogénea, por división basada en la característica estática y mediante regiones de Voronoi obtenidas por técnicas de agrupamiento y realizadas mediante politopos. Téngase en cuenta que al número de regiones mostrado no es indicativo de los méritos relativos de cada método. En general cabe esperar que los métodos más elaborados produzcan menor número de regiones proporcionando la misma capacidad de representación, pues eviten las zonas no visitadas del espacio $\mathcal{Z}$. En el ejemplo mostrado no se aprecia bien este hecho pues se trata de un caso con dimensión mínima.

- Forma de las regiones. En el algoritmo se ha considerado que las regiones son cajas en el espacio de puntos de funcionamiento. No existe ninguna restricción para considerar otro tipo de forma para las regiones, como elipsoides o politopos. Tales regiones vendrían descritas por $\|(z - z_p^f)^T M(z - z_p^f)\|_n < \Delta$ siendo n el índice de la norma utilizada $n \in [1, \infty)$ y M una matriz de transformación definida positiva.
- Funciones de ponderación. El valor de β afecta a la forma de la función de pertenencia. Interesa que las funciones de pertenencia correspondientes a regiones adyacentes tengan cierto grado de solape. Solamente de este modo es posible producir ponderaciones de los submodelos locales que por un lado tengan en cuenta el punto de operación y por otro lado consigan un cambio suave de las contribuciones de cada submodelo a medida que el punto de funcionamiento varía. Los resultados obtenidos por el algoritmo no cambian significativamente siempre y cuando los valores fijados para β no se acerquen a valores extremos que den como resultado regiones que no solapan o que solapan totalmente.
- Identificación. La identificación de los submodelos locales puede llevarse a cabo utilizando cualquiera de las técnicas que se emplean normalmente en aplicaciones industriales: curva de reacción ante escalón o impulso, mínimos cuadrados, o estadísticamente más elaborados como la identificación con regularizadores (Tikhonov and Arsenin, 1977). De hecho cada uno de los modelos locales puede diseñarse con criterios distintos y luego agruparse con los demás en la forma ya indicada.

Otra posibilidad consiste en la identificación simultánea de todos los parámetros de los modelos de forma similar a los sistemas neuroborrosos.

- Órdenes de los modelos locales. Estos órdenes pueden fijarse de varios modos: si se dispone de conocimiento a priori es conveniente utilizarlo, por ejemplo, en la industria de procesos es frecuente que la dinámica se pueda describir mediante un modelo de primer orden más retardo. En un caso general la mejor opción es realizar la selección de órdenes al mismo tiempo que la identificación usando para ello un criterio como el AIC de Akaike (Akaike, 1974). En cualquier caso es aconsejable fijar un valor máximo $omáx$ para los órdenes para evitar en lo posible problemas de sobreparametrización. Este límite se convierte entonces en un nuevo hiperparámetro del algoritmo.

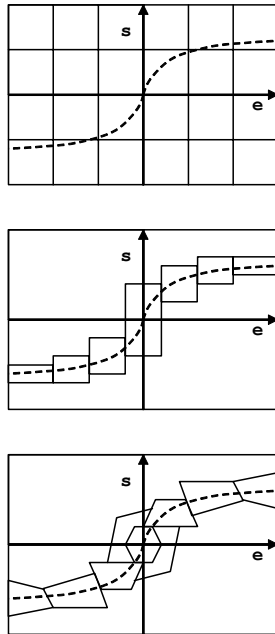


Figura 1. Resultados típicos al aplicar diversos métodos de selección de regiones a un sistema de una entrada y una salida.

3.3 Capacidad de representación de las AMLL

Resulta interesante comparar las capacidades representativas de las AMLL frente a otro tipo de modelos.

Modelos con no linealidades concentradas Los modelos de Hammerstein, de Wiener y de Volterra son casos particulares de estos modelos de no linealidades concentradas, los cuales se componen de bloques lineales unidos a no linealidades concentradas en determinados puntos del modelo (entrada, salida, realimentación) y es común que puedan caracterizarse por funciones lineales a trozos. Por ello conviene analizar cada tipo de no linealidad de forma separada.

Zona muerta La zona muerta se representa fácilmente mediante un punto de operación cuyo modelo asociado tiene ganancia nula. La amplitud de la zona muerta se hace corresponder con la extensión en el eje z_u de la región asociada al modelo. Surge el problema de calcular la extensión según el eje z_y de dicha región. Por un lado la salida en régimen permanente vale cero para valores de la entrada dentro de la zona muerta, este valor sin embargo conduciría a $\Delta_{y,p} = 0$ y a una región con volumen cero lo cual produciría una ponderación

nula. Es preciso modificar el algoritmo de construcción de modelos para tener en cuenta este caso de forma separada y proporcionar un valor adecuado para $\Delta_{y,p}$, por ejemplo $\Delta_{y,p} = 0,5(\max(y) - \min(y))$.

Saturación La saturación se representa también con facilidad si se escogen las regiones de modo que el valor de u que marca el comienzo de la saturación sea el punto de separación entre dos submodelos locales, uno de ellos previo a la saturación es un submodelo normal mientras que el otro tiene una región asociada infinita (no acotada por uno de los extremos) y representable por un modelo con ganancia nula. Es preciso nuevamente elegir un valor de $\Delta_{y,p} \neq 0$ para evitar que la región tenga un volumen nulo. Para tratar adecuadamente las regiones infinitas es preciso modificar la forma de las funciones de pertenencia de modo que contemplen condiciones asimétricas por encima o por debajo del punto de operación.

Cambios de pendiente Si la no linealidad se representa mediante una función lineal a trozos con cambios bruscos de pendiente es posible obtener un submodelo lineal local cuya región asociada coincida con cada uno de los trozos lineales. De este modo la agrupación de modelos locales puede representar con bastante exactitud modelos de este tipo. Es conveniente sin embargo modificar las funciones de pertenencia para que el cambio entre modelos sea abrupto, reflejando de este modo el cambio abrupto de pendiente. Para ello el solape entre regiones ha de tener un valor nulo lo cual se consigue con un valor de β infinito.

Transformaciones de Volterra Los modelos de Volterra usan no linealidades en la entrada (Gruber and Bordons, 2007). Es habitual que estas no linealidades sean funciones polinómicas como productos de variables, cuadrados, cubos, raíces cuadradas y sus combinaciones lineales. En ambos casos el tipo de no linealidad y su situación (entrada, salida, realimentación) viene sugerido por el conocimiento de los procesos que tienen lugar en la planta. Por ejemplo si se sabe que la salida y depende del cuadrado de la entrada entonces conviene introducir u^2 como variable de entrada del modelo.

Este tipo de modelos pueden reproducirse mediante AMLL utilizando submodelos lineales en los parámetros, como los representados por la ecuación (1), e incorporando en el vector regresor m variables que son valores pasados de la entrada y la salida y también productos de ellos como por ejemplo $y(k-1) \cdot u(k-1)$, $(u(k-3))^2$, etc.

Histéresis Para poder representar correctamente los sistemas con histéresis es preciso usar una variable w que indique cual de las varias ramas marca la evolución del sistema en un instante dado. Esta variable que depende de la historia pasada de la planta ha de ser identificada convenientemente mediante experimentos diseñados a tal efecto. La variación de esta variable puede expresarse mediante condiciones del tipo:

$$w(k) = \begin{cases} w(k-1), & \text{si } H(x(k-1)) = 0 \\ H(x(k-1)), & \text{en caso contrario} \end{cases}$$

siendo la función H discontinua dependiendo de un vector x que contiene un número de valores pasados de u y y adecuado. La histéresis más simple queda descrita de este modo utilizando la variable $x = y$ y la función

$$H(x) = \begin{cases} +1 & \text{si } x > y_{\max} \\ -1 & \text{si } x < y_{\min} \\ 0 & \text{en caso contrario} \end{cases}$$

Una vez obtenida, la variable w puede incorporarse al vector que define el punto de operación de modo que $z(k) = (u(k), y(k), w(k))$. Conviene además modificar las funciones de pertenencia para dar cabida a regiones infinitas. Para ello basta con considerar condiciones asimétricas por encima o por debajo del punto de operación.

Gracias a la inclusión de la variable w los modelos locales se pueden desarrollar de forma independiente para cada una de las ramas de la histéresis.

Aproximadores estáticos Las representaciones neuronales estáticas y las proporcionadas por los conjuntos borrosos son equivalentes y por tanto pueden considerarse como dos formas distintas de formular y usar una clase superior dada por series truncadas de funciones base. Desde este punto de vista existen semejanzas entre las AMLL y dichas series que conviene analizar. La diferencia fundamental estriba en que las AMLL combinan la salida de elementos dinámicos como son los submodelos locales, mientras que las representaciones neuronales y borrosas se obtienen como combinación de elementos estáticos actuando sobre el estado o sobre un vector de entrada que permite reconstruir la dinámica (Weigend and Gershenfeld, 1994).

Las series utilizan el mismo vector de entrada para realizar dos funciones diferentes: como regresor y como vector que define el punto de funcionamiento. Esta diferencia no es baladí pues el vector de entrada, como regresor que es, debe acomodar $na + nb$ señales. Este número es la dimensión del espacio que se divide en regiones lo cual conduce a un número potencial de regiones muy alto lo cual hace necesario introducir mecanismos simplificadores para la creación de funciones de pertenencia en los conjuntos borrosos y de selección de bases en las redes de neuronas.

Las ventajas de las AMLL son varias: en primer lugar una mayor sencillez de uso pues incorporan muy pocas modificaciones al caso del modelado lineal ejemplificado por el modelo de primer orden más tiempo muerto obtenido a partir de la curva de reacción. Como consecuencia de lo anterior las AMLL plantean menores requisitos en términos de cantidad de datos necesarios. La posibilidad de incorporar con facilidad conocimiento a priori sobre la planta es otra gran ventaja, conviene resaltar que este conocimiento a priori en un ambiente industrial puede referirse a ganancias, tiempos de respuesta, tipos de amortiguación de la respuesta, etc. que son fácilmente interpretables en términos de los parámetros de los MLL.

Además de estas ventajas obvias las AMLL sobresalen en un aspecto: su capacidad para utilizar submodelos de distintos órdenes y retardos para cada punto de funcionamiento. El paradigma neuronal/borroso en estos casos postula la necesidad de incluir como entradas potenciales del modelo todos los valores pasados que puedan ser necesarios para todos los puntos de funcionamiento. En el caso de plantas que exhiben cambios en la dinámica (retardos, tiempos característicos, etc.) se necesita que el vector de entrada del sistema neuronal o borroso tenga más de $na + nb$ variables. Para justificar esta afirmación baste pensar en un modelo que deba tener en cuenta un retardo variable entre 1 y 5 periodos de muestreo, las señales

a tener en cuenta en tal caso son $na + nb + 4$. El número de parámetros ajustables es proporcional al producto del número de entradas por el número de bases (Yuan and Fine, 1998), por lo que dicho número resulta ser también elevado y de ahí la necesidad de un gran número de datos para el entrenamiento. En el caso de las AMLL cada submodelo local puede ser tan simple como la situación permita, utilizando órdenes y retardos propios, de este modo la agrupación no se complica innecesariamente. Además cada submodelo puede identificarse de forma separada utilizando una cantidad de datos moderada.

La comparación no sería justa si no se indicasen algunas desventajas de las AMLL. Es por ello preciso indicar que en situaciones donde la abundancia de datos no es un problema, las AMLL producen resultados inferiores a los aproximadores neuronales y borrosos. Esto es debido a que los grados de libertad que ponen a disposición del usuario son usualmente y conscientemente desaprovechados: por un lado para intentar escapar del fenómeno de la sobreparametrización y por otro para facilitar la inclusión de conocimiento a priori y mantener cierta relación entre los parámetros ajustables y características de la planta como retardos, ganancias, etc. Los sistemas con aprendizaje por su parte prescinden de mantener dicha relación y optimizan durante el entrenamiento todos los parámetros ajustables de forma que se produzca un mejor acuerdo entre observaciones y las salidas del modelo en el conjunto de entrenamiento.

Redes recurrentes Las redes de neuronas artificiales que contienen lazos internos de realimentación son sistemas dinámicos en sí mismos. Por este motivo la capacidad para representar dinámica es excelente y no sufren los inconvenientes citados en el punto anterior para las redes acíclicas (estáticas). La gran flexibilidad que poseen para representar dinámicas diversas es sin duda superior a la de cualquiera de los métodos anteriormente descritos, incluyendo las AMLL. Pero no todo son ventajas, como inconvenientes cabe citar que el entrenamiento requiere una mayor cantidad de datos y que no es fácil incorporar conocimiento a priorístico como: valores conocidos de ganancias, tiempos de residencia, constantes de tiempos, etc.

3.4 Consideraciones prácticas

A la hora de utilizar el algoritmo de identificación propuesto es preciso tener en cuenta algunos factores que pueden tener gran influencia en el resultado.

Integradores. Si se sospecha que la dinámica que liga u e y contiene integradores puros en serie con el bloque que contiene la dinámica no lineal conviene eliminarlos en primer lugar. Para ello basta con tomar como salida la señal $\frac{d^v y}{dt}$ en lugar de y , siendo v el número de integradores puros colocados en serie.

Retardos puros. Los retardos puros son tenidos en cuenta por los modelos locales lineales, pero han de ser identificados previamente. Existe la posibilidad de usar el algoritmo repetidamente variando en cada prueba el retardo puro y seleccionando el valor que proporciona mejores resultados. Esta práctica no se recomienda pues puede dar lugar a que se obtengan modelos demasiado ajustados a unos datos concretos y con escasa capacidad de generalización.

Modelos multivariados. Los modelos multivariados pueden agruparse del mismo modo que los que constan de una sola entrada y una sola salida. La mayor complejidad proviene de la gestión de las regiones de operación cuyo número se

multiplica al incrementarse la dimensión de los espacios de entrada y/o salida. Para evitar que el número de regiones de operación sea excesivo se puede utilizar un algoritmo de agrupamiento. Este algoritmo proporciona los centros de las regiones a partir de datos tomados de la planta. Estos centros se sitúan cerca de los datos observados, de forma que las zonas sin datos no reciben ningún centro. De este modo se reduce bastante el número de regiones.

Selección de órdenes. El criterio de Akaike permite seleccionar automáticamente los órdenes de los MLL a partir de datos observados de la planta. Es conveniente sin embargo poner límite superior a los órdenes pues la selección hecha por el algoritmo que usa el criterio de Akaike depende de los datos suministrados. En ocasiones estos datos son poco variados por lo que el criterio se ve empujado a seleccionar órdenes altos para ajustarse a la realización particular del ruido presente en dichos datos.

Sobrep parametrización. Además de los órdenes de los MLL es preciso tener en cuenta el número de modelos que forman la agrupación. Un gran número de modelos puede dar como resultado una agrupación sobrep parametrizada cuya capacidad de generalización es pobre a pesar de que representa bien el comportamiento de la planta con los datos usados para identificar.

Excitación de las señales. El desarrollo de los modelos agregados conlleva una identificación en torno a un punto de trabajo para cada submodelo local. Esta identificación utiliza datos obtenidos, en el mejor de los casos, en bucle abierto con una señal suficientemente excitadora. El contenido en frecuencia de esta señal afecta a los resultados obtenidos en la identificación. De este hecho se derivan dos problemas potenciales: el primero consiste en garantizar que la excitación es suficiente para los submodelos locales, el segundo surge en el uso posterior de los modelos que puede producirse en situaciones distintas, posiblemente en bucle cerrado, por lo que se ven sometidos a señales de un contenido en frecuencia diferente.

Regularización. La regularización es una forma de introducir un sesgo en la identificación. Frecuentemente se utiliza como una forma de evitar la sobrep parametrización penalizando la complejidad de los modelos (Reed and Marks, 1998). En el contexto del control de procesos industriales la regularización puede también usarse para obtener modelos que proporcionan buenos resultados en la predicción multipaso o en simulación. Téngase en cuenta que la técnica habitual de identificación consiste en seleccionar los parámetros que minimizan el error a un paso. Este error en muchos casos es minimizado dando gran peso al término autorregresivo. Esto ocurre por ejemplo cuando los datos tienen poca variabilidad, o cuando el tiempo de muestreo es bajo respecto a los tiempos característicos de la planta. Como consecuencia los modelos resultantes pueden ser poco adecuados para la predicción multipaso. En tales casos la adición de términos regularizadores permite obtener submodelos más adecuados al uso posterior. Alternativamente se puede presentar este problema como un prefiltrado de los datos con relevancia para control (Rivera *et al.*, 1992), (Schwartz and Rivera, 2006).

Estas consideraciones han sido tenidas en cuenta a la hora de realizar las aplicaciones del método de creación de AMLL que se muestran a continuación.

4. CASOS PRÁCTICOS

A fin de ilustrar la aplicación del método propuesto y de analizar los resultados que produce se han seleccionado dos casos prácticos: un sistema simulado mediante una ecuación en diferencias y el sistema de captación de una planta real de producción de frío a partir de energía solar. Estos dos ejemplos prueban la posibilidad de usar el método como técnica simplificada aplicable en la industria.

4.1 Planta de Narendra

La llamada planta de Narendra ha sido usada como banco de pruebas de métodos de identificación y control desde su publicación en (Narendra and Parthasarathy, 1990). La dinámica es determinista y marcadamente alejada del caso lineal y viene dada por

$$y_{k+1} = \frac{y_k y_{k-1} y_{k-2} u_{k-1} (y_{k-2} - 1) + u_k}{1 + y_{k-1}^2 + y_{k-2}^2} \quad (7)$$

donde y representa la salida del sistema y u la entrada.

Como ilustración del comportamiento no lineal de esta planta véase la figura 2 que muestra la respuesta de la planta ante cambios en escalón en la entrada. Puede observarse que de un punto de funcionamiento a otro hay cambios en la ganancia estática, en la rapidez de la respuesta y en la amortiguación de la misma.

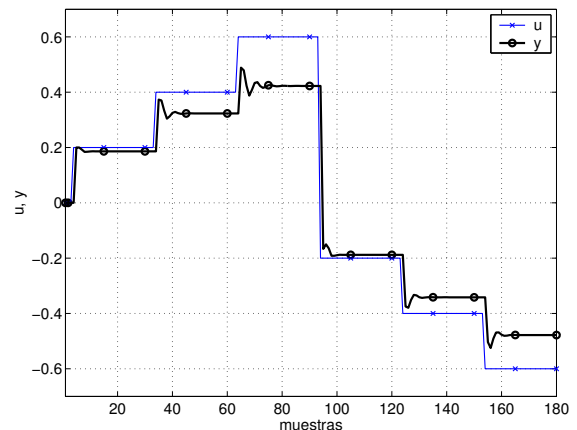


Figura 2. Entrada y salida de la planta de Narendra para un experimento en bucle abierto en el que se visitan distintos puntos de funcionamiento.

Señales de entrada Se utilizan señales del tipo $u(t) = A^{id} \sin(\omega^{id} t)$ una onda sinusoidal con una amplitud y frecuencia dada por la pareja (A^{id}, ω^{id}) . Un experimento de identificación consistirá en someter al sistema a varias señales de este tipo tomando las parejas (A^{id}, ω^{id}) mediante muestreo aleatorio con reemplazamiento del intervalo $[0, 1] \times [0, 20]$. De este modo se asegura el barrido en frecuencia y amplitud que permite cubrir las posibles zonas de trabajo.

Puntos de operación y regiones Para las pruebas se utiliza el algoritmo básico siendo el espacio de los puntos de funcionamiento es $z = (u, y)$. Se va a considerar un número np variable a fin de mostrar su efecto sobre el modelo obtenido. A partir de un np dado se empleará el método de selección de

regiones más simple posible consistente en la división a partir de la característica estática.

A fin de ilustrar el proceso, en la figura 3 se muestra la superficie de $\rho^{máx}(z) = \max \rho_i(z)$. Las mesetas que se observan corresponden a las regiones asociadas $\mathcal{R}_p$ de cada punto de operación. Dichos puntos coinciden con los centros de las mesetas y se sitúan sobre la característica estática de la planta.

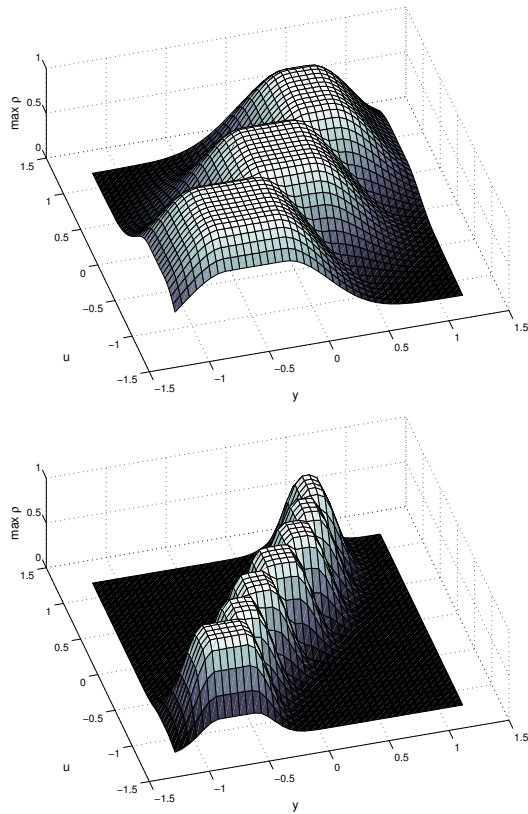


Figura 3. Superficie correspondientes a $\rho^{máx}(z)$ obtenidas por el algoritmo de creación de AMLL para $np = 3$ y $np = 7$ en la planta de Narendra.

Identificación Por un procedimiento simple de mínimos cuadrados y utilizando datos obtenidos de la simulación de la planta mediante la ecuación (7) se obtienen modelos locales para distintos puntos de operación. En cada caso los datos proporcionados se dividen en dos conjuntos: de entrenamiento (CE) y de prueba (CP). Se ha tomado $d = 0$ y np variable a fin de estudiar la influencia de este hiperparámetro. Los valores na y nb que definen los órdenes de los submodelos se han estimado a partir de los datos usando el criterio de Akaike, aunque se ha impuesto la limitación de un orden máximo ($omáx$).

Resultados Para cada pareja de hiperparámetros (np , $omáx$) el algoritmo ha dado como resultado una colección de MLL. A fin de poner a prueba la validez de la AMLL resultante se ha medido el error de generalización en la predicción a un paso y en simulación. Este error es la media cuadrática de los errores obtenidos en un conjunto de datos no usado previamente. La diferencia entre error de predicción y simulación es que en el primer caso se usan las medidas reales de la salida disponibles hasta $k - 1$ para producir la predicción para el instante k . En el caso de la simulación solamente se usan los valores iniciales

de la salida, de manera que las sucesivas predicciones se van apoyando en predicciones pasadas.

El conjunto de datos no utilizados previamente se ha seleccionado de forma que se visitan amplias zonas del espacio de trabajo y en particular zonas diferentes a las usadas en la identificación. En la gráfica superior de la figura 4 se muestra el error en la predicción a un paso. Puede verse que el menor error se obtiene para modelos generados por el algoritmo al utilizar $omáx=3$. En la parte inferior se muestra el error en simulación. En este caso el fenómeno de sobreparametrización se hace patente pues a partir de $np = 13$ las AMLL que usan modelos de orden mayor producen un error mayor.

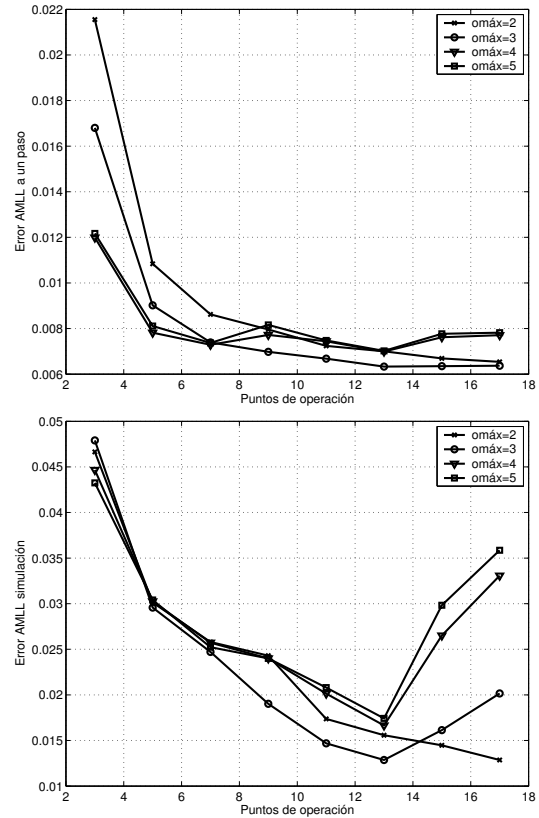


Figura 4. Error de generalización de las AMLL obtenidas con el algoritmo propuesto para la planta de Narendra en predicción a un paso (gráfica superior) y en simulación (gráfica inferior). Los valores de los hiperparámetros se indican en el eje horizontal (np) y en las distintas curvas ($omáx$).

Puede observarse la mejor adecuación de los modelos locales para valores altos de $np = 5$ respecto a valores menores, como era de esperar debido a que la región en la que operan es menor cuanto mayor es el número de modelos locales.

4.2 Aplicación a una planta de frío solar

La planta de producción de frío mediante energía solar se encuentra situada en los laboratorios de la Escuela Superior de Ingenieros de la Universidad de Sevilla. Actualmente está siendo utilizada como banco de pruebas para investigaciones sobre control híbrido dentro del marco de la Red de Excelencia HyCon. Además ha sido objeto de estudios en diferentes investigaciones: (Núñez-Reyes *et al.*, 2005; Núñez-Reyes, 2007; Zambrano, 2007).

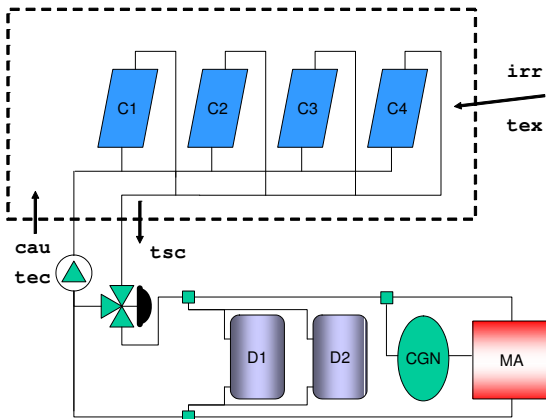


Figura 5. Fotografía y diagrama simplificado de la planta de producción de frío mediante energía solar incluyendo los captadores (C1 a c4), los depósitos auxiliares (D1 y D2) y la máquina de absorción (MA). En el recuadro de línea de trazos se ha enmarcado el circuito de captación.

La planta solar se puede analizar como una instalación de aire acondicionado que usa energía térmica para producir aire frío. Los principales componentes son los siguientes:

- Sistema solar o de captación de baja temperatura, compuesto por un conjunto de colectores solares (C1 a C4), con una superficie de captación de 151 m^2 y una potencia nominal de 50 kW .
- Sistema de almacenamiento, compuesto por dos depósitos acumuladores (D1, D2) de $2,500 \text{ m}^3$ cada uno interconectados entre sí.
- Sistema de aporte auxiliar de energía, consistente en una caldera de gas natural (CGN).
- Sistema de refrigeración, compuesto por una máquina de absorción (MA) de BrLi con una potencia de 35 kW .

La planta solar permite diferentes configuraciones en las que los sistemas que la componen pueden funcionar de manera independiente. A efectos de modelado puede dividirse en varios subsistemas: captación, almacenamiento, aporte auxiliar de energía, máquina de absorción y carga. Cada subsistema puede modelarse separadamente pues se dispone de medidas de muchas variables intermedias: temperaturas y caudales fundamentalmente. De esta forma cada subsistema es tratado como un módulo independiente cuyo modelo puede ser creado y validado con mayor facilidad. Para las pruebas de las AMLL se va a considerar el sistema de captación de energía solar.

Es de resaltar que la construcción de modelos para este tipo de plantas tiene un gran interés pues mediante un control adecuado

se pueden lograr mejoras en el rendimiento energético de las mismas (Rubio *et al.*, 2006). Una de las estrategias de control que mejor se adaptan a esta situación es el control predictivo (Camacho and Bordons, 2004; Francisco and Vega, 2006; García-Nieto *et al.*, 2008) pues permite optimizar las señales de control o las consignas para los controladores locales (Cirre *et al.*, 2004) de forma que se mejore el rendimiento final teniendo en cuenta al mismo tiempo restricciones de operación.

Descripción La figura 5 muestra un diagrama simplificado de la planta de producción de frío. La zona dentro del recuadro es el circuito de captación de calor, para el cual se han indicado las variables de interés:

- Caudal de agua que circula por los colectores en l/h . Se utilizará un valor normalizado obtenido dividiendo por 8000. Se emplea la abreviatura cau.
- Irradiación solar en kW/m^2 . Se utilizará un valor normalizado obtenido dividiendo por 1000. Se emplea la abreviatura irr.
- Temperatura exterior en $^{\circ}\text{C}$. Se utilizará un valor normalizado obtenido dividiendo por 100. Para nominar a este valor normalizado se utilizará la abreviatura tex.
- Temperatura de entrada de colectores en $^{\circ}\text{C}$. Nuevamente se normaliza dividiendo por 100, empleándose la abreviatura tec.
- Temperatura de salida de colectores en $^{\circ}\text{C}$. También se divide por 100 para normalizar y la abreviatura empleada es tsc.
- Incremento de temperatura en los colectores en $^{\circ}\text{C}$. Esta variable no se mide directamente sino que ha de calcularse a partir de las temperaturas de entrada y salida de colectores. Se usará un valor normalizado obtenido como $itc(k) = 3(tsc(k) - tec(k - d))$ siendo d el número de periodos de muestreo que tarda el agua en recorrer los colectores para un caudal medio.

Datos Los datos consisten en la evolución de las variables de interés tomadas en ensayos realizados durante la operación normal de la planta más algún otro realizado a propósito para la identificación. Estos ensayos incluyen condiciones de operación diversas: en lazo abierto y cerrado y con la planta funcionando en distintos modos: recirculación, producción directa, producción con depósito, etc. Los datos han sido normalizados, tratados para eliminar falsas medidas y remuestreados a 20s. La figura 6 muestra la evolución típica de las variables en un día sin nubes correspondiente al 20 de Marzo de 2006.

En total se ha contado con datos correspondientes a 42 días de Marzo a Octubre de 2006. Los días se han repartido de manera aproximadamente uniforme para dar lugar a tres conjuntos de datos habitualmente usados en identificación. El conjunto de entrenamiento (CE) se utiliza para adaptar los modelos. El conjunto de prueba (CP) se utiliza para elegir los hiperparámetros. Finalmente el conjunto de validación (CV) no es usado en modo alguno durante la construcción de modelos y sirve para calcular la capacidad de generalización de la AMLL resultante.

Puntos de operación y regiones En primer lugar se procede a estimar los puntos de régimen permanente. Esta tarea no puede realizarse con experimentos de manera fácil pues algunas variables no son manipulables como la irradiación. En lugar de eso se ha utilizado un método para hallar puntos de régimen estacionario o cuasi-estacionario a partir de experimentos realizados en distintas condiciones de operación. El método consiste

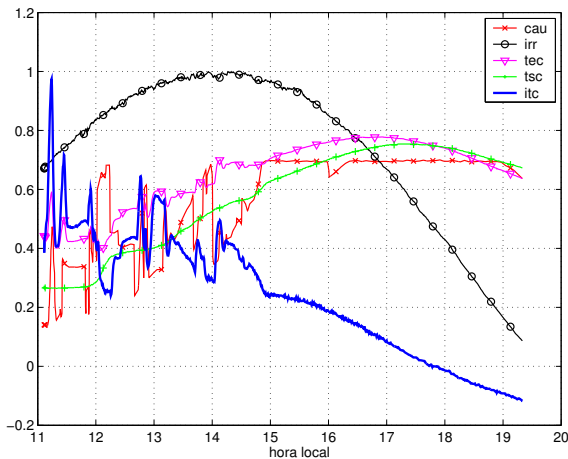


Figura 6. Evolución típica de las variables del captador de la planta de frío solar. Se representan datos normalizados tomados el 20 de Marzo de 2006.

simplemente en buscar en los archivos intervalos temporales donde las variables apenas modifican sus valores.

A partir de los valores de régimen permanente se eligen los puntos de funcionamiento que dan lugar a los MLL. El espacio considerado es $z = (cau, irr, itc)$. El método de selección de regiones consiste en la división homogénea del intervalo de variación de las señales de entrada, colocando un punto de funcionamiento únicamente en las regiones que contengan datos. A modo de ejemplo se incluye la figura 7 que ha sido realizada para un número de submodelos $np = 17$.

Identificación La identificación se lleva a cabo varias veces variando los hiperparámetros del algoritmo para de este modo poder extraer resultados comparativos. En todos los casos se ha utilizado solamente el CE para la identificación y el CP para dirigir la búsqueda de los mejores hiperparámetros.

Antes de proceder a la identificación ha sido necesario calcular unos valores adecuados para los retardos puros para cada sub-modelo local. Se ha elegido para ello un procedimiento simple consistente en la inspección visual de los datos disponibles. En algunos de los experimentos se han hallado cambios bruscos en variables de entrada que producen un efecto en la salida cierto tiempo después. De este modo se ha llegado a establecer unos retardos que dependen fundamentalmente del caudal.

Los órdenes na , nb y nc de los submodelos locales han sido obtenidos a partir de los datos, pero fijando un valor máximo $omáx$ a fin de poder realizar comparaciones.

Resultados En primer lugar se presentan los resultados obtenidos en predicción a un paso y en simulación para el caso $np = 17$ y $omáx = 2$, tomando los datos de un día perteneciente al CV. En la figura 8 se muestra, en la parte superior la evolución de las entradas, en la parte central la salida real y la salida dada por la AMLL, finalmente, en la parte inferior se ha dibujado la evolución de los errores. Puede verse que el error en simulación de la AMLL raras veces supera los 5 grados Celsius y cuando lo hace es de forma puntual. Téngase en cuenta que el uso de estos modelos no suele ser la simulación sino la predicción en un horizonte de 1 a varias decenas de instantes de muestreo. En tales casos el error es bastante menor y totalmente adecuado para los requisitos de control de la planta de producción de frío.

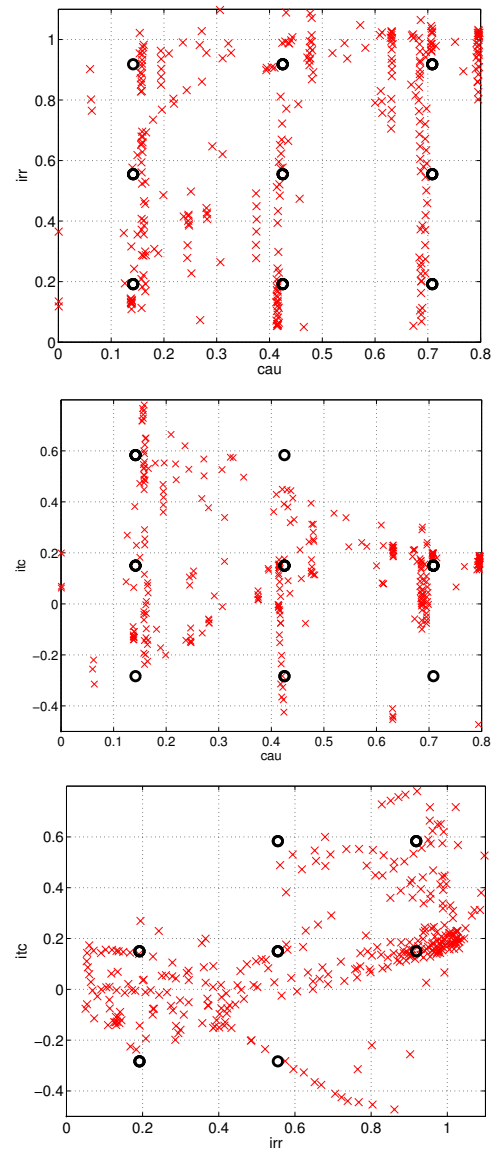


Figura 7. Puntos de funcionamiento seleccionados (círculos) y algunos valores medidos en régimen cuasi-estacionario sobre la planta (cruces) mostrados en las tres proyecciones bidimensionales del espacio (cau, irr, itc).

Tabla 1 Error cuadrático medio de simulación obtenido por las AMLL con np y $omáx$ fijados con el CP

np	$omáx$	ECM_{CE}	ECM_{CP}	ECM_{CV}
12	3	2.5241	2.8670	2.8443
17	3	2.1563	2.5730	2.6918
26	5	2.1464	2.6056	2.6466
36	1	2.3732	2.7267	2.9941

A fin de comprobar la influencia de los hiperparámetros en el resultado de la identificación se han realizado pruebas exhaustivas variando el número de submodelos np y del orden máximo $omáx$. En ellas se han dividido los datos disponibles en particiones de 200 muestras. Sobre cada partición se ha medido el error en simulación proporcionado por la AMLL. Finalmente se ha calculado el error cuadrático medio para todas las particiones. En la tabla 1 se muestran los resultados obtenidos.

Puede observarse que incluso para un número de puntos de funcionamiento moderado se obtienen resultados adecuados para el control de la planta pues una desviación de dos o tres

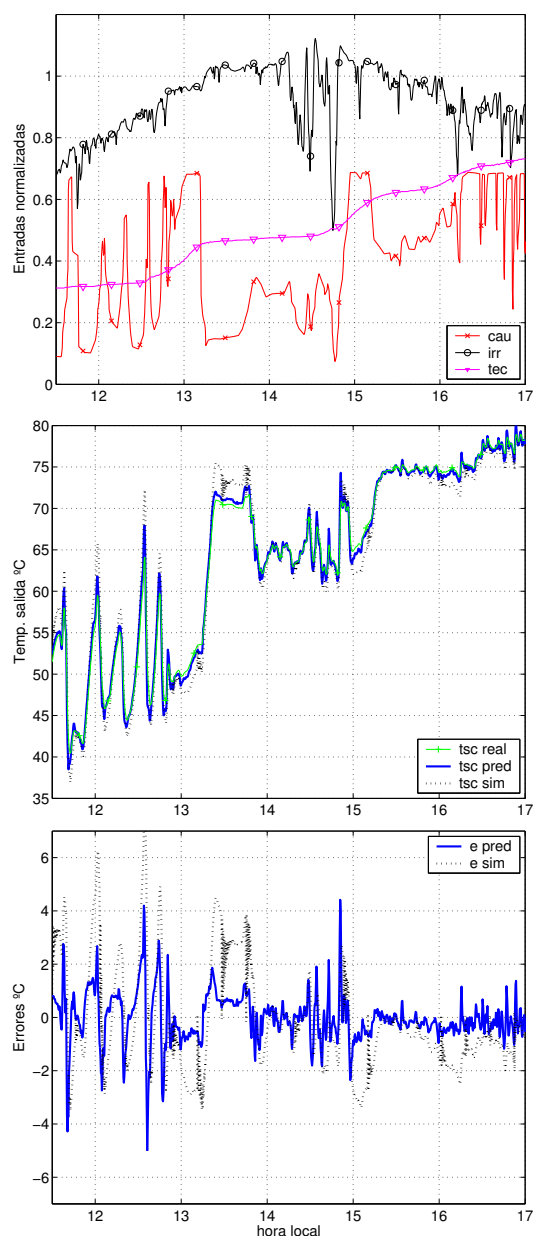


Figura 8. Resultados obtenidos en predicción a un paso y en simulación en un día del CV.

grados en la predicción a largo plazo no es significativa. La capacidad de generalización queda también patente al observar que los errores en CV no son muy diferentes a los obtenidos para el CP.

5. CONCLUSIONES

El algoritmo propuesto para la creación de AMLL utiliza exclusivamente señales instantáneas de entrada y salida del sistema. Se ha mostrado que esta elección simplifica algunos aspectos de la creación del modelo, manteniendo intacta la capacidad de representación. El análisis realizado en cuanto a capacidad de representación revela que las AMLL pueden usarse para modelar plantas como otros métodos como las series de Volterra o los modelos de Hammerstein. La comparación realizada con las redes de neuronas artificiales y los modelos basados en conjuntos borrosos permite discernir en qué situaciones pueden esperarse mejores resultados por uno u otro método. En particular se ha

argumentado que para la práctica industrial las AMLL plantean menores requisitos prácticos al tiempo que facilitan la inclusión de conocimiento apriorístico.

Finalmente el método propuesto ha sido probado en dos escenarios: una planta simulada monovariable y una planta real con varias entradas. En este segundo caso se ha presentado la dificultad añadida de tratarse de una planta con retrasos variables en función del punto de operación. Los resultados mostrados indican que las AMLL son una posibilidad a tener en cuenta.

AGRADECIMIENTOS

Este trabajo ha sido financiado parcialmente por la Comisión de la Comunidad Europea a través del proyecto de investigación HYCON FP6-511368.

REFERENCIAS

- Akaike, H. (1974). A new look at the statistical model identification. *IEEE Trans. on Automatic Control* **AC-19**(5), 716–723.
- Alamo, T., A. Cepeda, D. Limon and E. F. Camacho (2006). A new concept of invariance for saturated systems. *Automatica* **42**(9), 1515–1521.
- Bellman, R. E. (1957). *Dynamic Programming*. Princeton University Press, Princeton.
- Bontempi, G. and M. Birattari (2005). From linearization to lazy learning: A survey of divide-and-conquer techniques for nonlinear control. *International Journal of Computational Cognition* **3**(1), 56–73.
- Camacho, E. F. and C. Bordons (2004). Control predictivo: Pasado, presente y futuro. *Revista Iberoamericana de Automática e Informática Industrial* **1**(3), 5–28.
- Cirre, Cristina M., Loreto Valenzuela, Manuel Berenguel and Eduardo Camacho (2004). Control de plantas solares con generación automática de consignas. *Revista Iberoamericana de Automática e Informática Industrial* **1**(1), 50–56.
- de Prada, C. (2004). El futuro del control de procesos. *Revista Iberoamericana de Automática e Informática Industrial* **1**(1), 5–14.
- Foss, B. A., T. A. Johansen and A. V. Sorenson (1995). Nonlinear predictive control using local models - applied to a batch fermentation process. *Control Engineering Practice* **3**(3), 389–396.
- Francisco, M. and P. Vega (2006). Diseño integrado de procesos de depuración de aguas utilizando control predictivo basado en modelos. *Revista Iberoamericana de Automática e Informática Industrial* **3**(4), 88–98.
- Fritzke, B. (1994). Growing cell structures - a self-organizing network for unsupervised and supervised learning. *Neural Networks* **7**(9), 1441–1460.
- García-Nieto, S., M. Martínez, X. Blasco and J. Sanchis (2008). Nonlinear predictive control based on local model networks for air management in diesel engines. *Control Engineering Practice*.
- Ghraizi, R. A., E. Martínez, C. de Prada, F. Cifuentes and J.L. Martínez (2007). Performance monitoring of industrial controllers based on the predictability of controller behavior. *Computers and Chemical Engineering* **31**, 477–486.
- Gruber, J. K. and C. Bordons (2007). Control predictivo no lineal basado en modelos de volterra. aplicación a una planta piloto. *Revista Iberoamericana de Automática e Informática Industrial* **4**(3), 56–73.

- Hunt, K. J., T. A. Johansen, J. Kalkkuhl, H. Fritz and Th. Gottsche (2000). Nonlinear speed control for an experimental vehicle using a generalized gain scheduling approach. *IEEE Transactions on Control Systems Technology* **8**(3), 381–395.
- Johansen, T. A. and B. A. Foss (1993). Constructing narmax models using armx models. *International Journal of Control* **58**(5), 1125–1153.
- Johansen, T. A. and B. A. Foss (1995). Identification of nonlinear system structure and parameters using regime decomposition. *Automatica* **31**(2), 321–326.
- Johansen, T. A. and B. A. Foss (1998). ORBIT—operating regime based modelling and identification toolkit. *Control Engineering Practice* **6**(10), 1277–1286.
- Johansen, T. A., K. J. Hunt and I. Petersen (2000). Gain-scheduled control of a solar power plant. *Control Engineering Practice* **8**(9), 1011–1022.
- Moody, J. and C. Darken (1989). Fast learning in networks of locally-tuned processing units. *Neural Computation* **1**, 281–294.
- Murray, R. and T. Johansen (1997). *Multiple Model Approaches to Modelling and Control*. Taylor and Francis. UK.
- Narendra, K. S. and K. Parthasarathy (1990). Identification and control of dynamical systems using neural networks. *IEEE Transactions on Neural Networks* **1**, 4–27.
- Núñez-Reyes, A. (2007). Contribuciones al Control de Plantas de Producción de Frío Mediante Energía Solar. PhD thesis. Departamento de Ingeniería de Sistemas y Automática. Escuela Superior de Ingenieros. Universidad de Sevilla.
- Núñez-Reyes, A., J.E. Normey-Rico, C. Bordons and E.F. Camacho (2005). A smith predictive based MPC in a solar air conditioning plant. *Journal of Process control* **15**–1, 1–10.
- Principe, J. C., L. Wang and M. A. Motter (1998). Local dynamic modeling with self-organizing maps and applications to nonlinear system identification and control. *Proceedings of the IEEE* **86**(11), 2240–2257.
- Reed, R. D. and R. J. Marks (1998). *Neural smithing*. MIT press, Cambridge.
- Rivera, D. E. (2007). Una metodología para la identificación integrada con el diseño de controladores imc-pid. *Revista Iberoamericana de Automática e Informática Industrial* **4**(4), 5–18.
- Rivera, D. E., J. F. Pollard and C. E. García (1992). Control-relevant prefiltering: a systematic design approach and case study. *IEEE Transactions on automatic control* **37**(7), 964–974.
- Rodriguez, M., D. Sarabia and C. de Prada (2007). Hybrid predictive control of a simulated chemical plant. In: *Taming Heterogeneity and Complexity of Embedded Control* (F. Lamnabhi-Lagarigue, S. Laghrouche, A. Loria and E. Panteley, Eds.). pp. 617–634. ISTE.
- Rubio, F. R., E. F. Camacho and M. Berenguel (2006). Control de campos de colectores solares. *Revista Iberoamericana de Automática e Informática Industrial* **3**(4), 26–45.
- Sarabia, D., C. De Prada, S. Cristea and R. Mazaeda (2006). Hybrid model predictive control of a sugar end section. *Computer Aided Chemical Engineering* **21**–12, 1269–1274.
- Schwartz, J. D. and D. E. Rivera (2006). Control-relevant demand modeling for supply chain management. In: *Proceedings of the 14th Annual IFAC Symposium on System Identification*. Newcastle, Australia. pp. 267–272.
- Tikhonov, A.Ñ. and V. Y. Arsenin (1977). *Solutions of ill-posed problems*. V.H. Winston and sons, Washington.
- Weigend, A.S. and N.A. Gershenfeld (1994). *Time Series Prediction*. Addison-Wesley.
- Yuan, J.L. and T.L. Fine (1998). Neural network design for small training sets of high dimension. *IEEE Transactions on Neural Networks* **9**, 266–380.
- Zambrano, D. (2007). Modelado y control predictivo híbrido de una planta de refrigeración solar. PhD thesis. Departamento de Ingeniería de Sistemas y Automática. Escuela Superior de Ingenieros. Universidad de Sevilla.