

FRANCISCO ANTONIO LOTUFO

DESENVOLVIMENTO DE UM SENSOR VIRTUAL PARA PROCESSOS
NÃO-LINEARES E VARIANTES NO TEMPO, COM APLICAÇÃO
EM PLANTA DE NEUTRALIZAÇÃO DE pH

Tese apresentada à Faculdade de Engenharia
do *Campus* de Guaratinguetá, Universidade
Estadual Paulista, para a obtenção do título
de Doutor em Engenharia Mecânica na área
de Projetos.

Orientador: Prof. Dr. Samuel E. de Lucena

Guaratinguetá
2010

L884d	Lotufo, Francisco Antonio Desenvolvimento de um sensor virtual para processos não lineares e variantes no tempo, com aplicação em planta de neutralização de pH / Francisco Antonio Lotufo. - Guaratinguetá : [s.n.], 2010 131 f.: il. Bibliografia: f. 118-126 Tese (Doutorado) – Universidade Estadual Paulista, Faculdade de Engenharia de Guaratinguetá, 2010 Orientador: Prof. Dr. Samuel Euzédice de Lucena 1. Identificação de sistemas 2. Lógica difusa 3. Detectores I. Título CDU 681.5.015
-------	--



UNESP **UNIVERSIDADE ESTADUAL PAULISTA**
Faculdade de Engenharia do Campus de Guaratinguetá

**"DESENVOLVIMENTO DE UM SENSOR VIRTUAL PARA
PROCESSOS NÃO-LINEARES E VARIANTES NO TEMPO, COM
APLICAÇÃO EM PLANTA DE NEUTRALIZAÇÃO DE pH"**

FRANCISCO ANTONIO LOTUFO

ESTA TESE FOI JULGADA ADEQUADA PARA A OBTENÇÃO DO TÍTULO DE
"DOUTOR EM ENGENHARIA MECÂNICA"

PROGRAMA: ENGENHARIA MECÂNICA
ÁREA: PROJETOS

APROVADA EM SUA FORMA FINAL PELO PROGRAMA DE PÓS-GRADUAÇÃO

Prof. Dr. Marcelo dos Santos Pereira
Coordenador

BANCA EXAMINADORA:

Prof. Dr. SAMUEL E. DE LUCENA
Orientador/UNESP-FEG

Prof. Dr. INÁCIO BIANCHI
UNESP-FEG

Prof. Dr. LEONARDO MESQUITA
UNESP-FEG

Prof. Dr. GERMANO LAMBERT-TORRES
UNIFEI – Universidade Federal de Itajubá

Prof. Dr^a. NEUSA MARIA FRANCO DE OLIVEIRA
ITA – Instituto Tecnológico de Aeronáutica

Outubro de 2010

DADOS CURRICULARES

FRANCISCO ANTONIO LOTUFO

NASCIMENTO	09.10.1963 – TAUBATÉ / SP
FILIAÇÃO	Luiz de Gonzaga Lotufo Olarpha Garcez Lotufo
1983/1987	Curso de Graduação Engenharia Elétrica - Universidade de Taubaté
1997/1998	Curso de Pós-Graduação em Engenharia Elétrica, nível de Mestrado, na Universidade Federal de Itajubá.
2009/2010	Curso de Pós-Graduação em Engenharia Mecânica, nível de Doutorado, na Faculdade de Engenharia do <i>Campus</i> de Guaratinguetá da Universidade Estadual Paulista.

de modo especial, à minha esposa, Ana Lucia, e à minha filha, Ana Beatriz, que foram as grandes motivadoras para que eu continuasse e terminasse o curso.

AGRADECIMENTOS

Em primeiro lugar agradeço a Deus, pela minha vida, minha inteligência, minha família e meus amigos,

ao meu orientador, *Prof. Dr. Samuel E. de Lucena*, pelo seu incondicional apoio e incentivo. Sem a sua orientação, dedicação e auxílio, o estudo aqui apresentado seria praticamente impossível.

aos meus pais, *Luiz (in memoriam) e Olarpha (in memoriam)*, que, apesar das dificuldades, souberam orientar, educar e incentivar a mim e a meus irmãos.

aos meus colegas de Departamento, que sempre me acolheram, respeitaram e apoiaram em todos esses anos de fraterno convívio,

às funcionárias da Biblioteca do *Campus* de Guaratinguetá, e em especial à colega *Ana Maria*, pela dedicação, presteza e principalmente pela vontade de ajudar.

“I often say that when you can measure what you are speaking about, and express it in numbers, you know something about it; but when you can not express it in numbers, your knowledge is of a meagre and unsatisfactory kind; it may be the beginning of knowledge, but you have scarcely, in your thoughts, advanced to the state of science, whatever the matter may be.”

William Thomson, Lord Kelvin

“Há três maneiras de agir sabiamente:
A primeira pela meditação, que é a mais sábia
a segunda pela imitação, que é a mais fácil
a terceira pela experiência, que é a mais amarga”

Confúcio

LOTUFO, F. A. **Desenvolvimento de um sensor virtual para processos não-lineares e variantes no tempo, com aplicação em planta de neutralização de pH.** 2010. 131 f. Tese (Doutorado em Engenharia Mecânica) – Faculdade de Engenharia do *Campus* de Guaratinguetá, Universidade Estadual Paulista, Guaratinguetá, 2010.

RESUMO

Este trabalho apresenta uma metodologia para desenvolvimento de sensores virtuais capazes de inferir variáveis de processos altamente não lineares e variantes no tempo. A metodologia proposta emprega modelagem nebulosa de sistemas dinâmicos complexos, em que a parte antecedente das regras emprega a técnica de agrupamento (clusterização) de Gustafson-Kessel (*product space clustering*), e os parâmetros da parte conseqüente das regras são estimados utilizando-se o algoritmo dos mínimos quadrados recursivos, com fator de esquecimento variável. O algoritmo proposto foi avaliado por meio de um modelo virtual implementado em ambiente Matlab®/Simulink® para o processo de neutralização de pH, amplamente utilizado pela literatura técnico-científica para investigação de algoritmos de identificação, controle de processos e simulação de sistemas não lineares e variantes no tempo. O algoritmo de identificação nebulosa proposto e implementado neste trabalho, utilizado como sensor virtual de pH, forneceu resultados muito coerentes, quando comparado com outras técnica de modelagem da literatura, no tocante ao tempo de resposta, erro de predição, capacidade de adaptação e número de amostras necessárias à fase de treinamento.

PALAVRAS-CHAVE: Sensor Virtual. Sensor Inteligente. Identificação de Sistema. Lógica Nebulosa. Agrupamento Nebuloso.

LOTUFO, F. A. **Development of a virtual sensor for nonlinear and time-varying processes with application to pH neutralization plant.** 2010. 131 f. Tese (Doutorado em Engenharia Mecânica) – Faculdade de Engenharia do Campus de Guaratinguetá, Universidade Estadual Paulista, Guaratinguetá, 2010.

ABSTRACT

The design of a virtual sensor involves the choice of variables to be measured and the choice of a method for obtaining the model. This paper presents a methodology for developing virtual sensors capable of inferring process variables highly nonlinear and time varying. The proposed methodology employs fuzzy modeling of complex dynamic systems in which the antecedent part of rules employs the technique of clustering Gustafson-Kessel (product space clustering), and the parameters of consequent part of rules are estimated using the algorithm of the minimum recursive squares with variable forgetting factor. The algorithm was evaluated through a virtual model implemented in Matlab®/Simulink® for the pH neutralization process, widely used by scientific and technical literature to investigate the identification algorithms, process control and simulation of nonlinear systems and variants in time. The fuzzy identification algorithm proposed and implemented in this work, used as a virtual sensor of pH, provided very interesting results when compared with other modeling technique of the literature regarding the response time, error prediction, adaptive capacity and number of samples necessary for the training phase.

KEYWORDS: Virtual Sensors. Intelligent Sensors. Systems Identification. Fuzzy Logic. Fuzzy Clustering.

LISTA DE FIGURAS

FIGURA 1 – Ilustração da escala do pH e alguns exemplos com substâncias comuns.	23
FIGURA 2 – Ilustração dos eletrodos de vidro para medição de pH.	25
FIGURA 3 – Curvas de titulação para: a) ácido forte – base forte, b) ácido fraco – base forte, c) ácido forte – base fraca, d) ácido fraco – base fraca	27
FIGURA 4 - Curvas de titulação para soluções <i>buffered</i> e <i>unbuffered</i>	29
FIGURA 5 – Cadeia de medição analógica.....	31
FIGURA 6 – Cadeia de medição digital.....	31
FIGURA 7 – Sensor virtual conectado a uma planta	34
FIGURA 8 – Os sensores virtuais SV1 e SV2 são usados em duas malhas de controle quando o sensor XT é compartilhado no tempo entre as duas plantas.....	35
FIGURA 9 – Sensor virtual substituindo um sensor real inexistente.....	35
FIGURA 10 – Sensor virtual usando um modelo de planta quando nenhum sensor real está disponível	36
FIGURA 11 – Abordagem de identificação baseada em agrupamento nebuloso	44
FIGURA 12 - Interpretação baseada em regra de agrupamentos nebulosos.....	56
FIGURA 13 - Agrupamentos nebulosos hiper-elipsoidal	56
FIGURA 14 – a- Ilustração de agrupamentos gerados pelo FCM. b- Ilustração de agrupamentos gerados por Gustafson-Kessel.....	61
FIGURA 15 – Ilustração das regiões do espaço de entrada associadas às regras.	67
FIGURA 16 – Exemplo do espaço de entradas para o modelo de Wang-Langari.....	68
FIGURA 17 – Diferenças do resultado da clusterização. A primeira utiliza apenas a variável de entrada, a segunda utiliza ambas as variáveis.....	69
FIGURA 18 – Ilustração do espaço de entrada para o modelo “Regras Limitadas”	72
FIGURA 19 – Uma ilustração de um processo de neutralização de pH	79
FIGURA 20 – Sistema de controle típico de pH	80
FIGURA 21 – Curva de titulação para ácido fraco neutralizado por base forte	83
FIGURA 22 – Curva de titulação para ácido forte neutralizado por base forte	84
FIGURA 23 – Diagrama geral do processo de neutralização de pH.....	85
FIGURA 24 – Resposta de malha aberta para o teste 1: (a) Vazão de ácido; (b) Concentração de ácido; (c) Vazão de <i>buffer</i> ; (d) pH do fluxo de saída.....	95

FIGURA 25 – Resposta de malha aberta para o teste 2: (a) vazão de ácido; (b) concentração de ácido; (c) vazão de <i>buffer</i> ; (d) pH do fluxo de saída	96
FIGURA 26 – Sinais de entrada (a) e saída (b) da Equação (63) com variação lenta dos parâmetros.....	99
FIGURA 27 – (a) Parâmetros a e b , real (---) e estimado (---). (b) imagem retirada de Moustafa (1983)...	100
FIGURA 28 – (a) Fator de esquecimento variável, (b) Erro de saída entre o valor simulado e o estimado.....	101
FIGURA 29 – Sinais de entrada (a) e saída (b) da Equação (63) com variação abrupta dos parâmetros	102
FIGURA 30 – Parâmetros a e b , real (---) e estimado (---).....	103
FIGURA 31 – (a) Fator de esquecimento variável, (b) Erro de saída entre o valor simulado e o estimado.....	103
FIGURA 32 – (a) Sinal de entrada aleatório, (b) Sinal de entrada periódico	105
FIGURA 33 – (a) Sinais de saída real (---) estimado (---) - treinamento, (b) Sinais de saída real (---) estimado (---) - validação.....	106
FIGURA 34 – Valor RMS do erro e valor do erro de predição.	106
FIGURA 35 – Sinais de saída real (---) estimado (---) - validação.	108
FIGURA 36 – Modelo simplificado do tanque de neutralização de pH utilizado.	108
FIGURA 37 – (a) Esquema de simulação usado no treinamento, (b) Esquema usado como sensor virtual	109
FIGURA 38 – (a) Entrada vazão de base q_3 , (b) saída pH medido.	110
FIGURA 39 – (a) Sinais de saída real (•••) e estimado (---) – treinamento.	111
FIGURA 40 – (a) Sinais de saída real (•••) e estimado (---) – validação, (b) figura que representa o modelo de validação em comparação com o valor simulado extraído de Babuška e Verbruggen (1997).....	112
FIGURA 41 – Valor do erro de predição.	113

LISTA DE TABELAS

TABELA 1 – Informação <i>crisp</i> e nebulosa (<i>fuzzy</i>) em sistemas.....	47
TABELA 2 – Aplicação do pH em produções industriais	79
TABELA 3 – Instalação industrial e bancada de laboratório sob as condições normais de estado estável NSS.....	93
TABELA 4 – Teste 1: degraus de q_1 com diferentes valores de <i>buffer</i>	95
TABELA 5 – Teste 2: onda senoidal com período de 15 min W_{a1} com/sem <i>buffer</i>	96
TABELA 6 – Variação do erro RMS em função do número de amostras de treinamento	107

LISTA DE SÍMBOLOS

$y_f(k)$	função no instante k	
N_p	regressões da p-ésima entrada	
N_f	regressões da saída	
pH	potencial hidrogeniônico	
$\mathbf{R}^n$	espaço multidimensional n	
$\mathbf{u}$	entradas	
$\mathbf{y}$	saídas	
f_{NL}	função não-linear	
θ	parâmetros	
$[H^+]$	concentração dos íons de hidrogênio	M
OH^-	íons hidroxila	mol
E_{pH}	tensão através da membrana de vidro	mV
R	constante universal dos gases	$J.K^{-1}.mol^{-1}$
T	temperatura	K
F	constante de Faraday	$C.mol^{-1}$
A_i	termos linguísticos antecedentes	
B_i	termos linguísticos consequentes	
$\mu(x)$	função de pertinência dos conjuntos nebulosos antecedentes	
v_i	centro do agrupamento i	
μ_{ik}	pertinência do ponto k ao agrupamento i	
c	número desejado de agrupamentos	
$\mathbf{B}$	vetor de parâmetros	
λ_k	fator de esquecimento	
W_a	invariante de reação	
θ_{pH}	tempo morto de medição do pH	s
θ_{mix}	atraso na mistura	s
q_1	fluxo de ácido - influente	m^3/s
q_4	fluxo de saída – efluente	m^3/s
$[C_3]$	concentração do fluxo de base -[NaOH]	M

SUMÁRIO

1 INTRODUÇÃO	15
2 SENSORES VIRTUAIS DE pH.....	22
2.1 Definição de pH.....	22
2.2 Sensor de pH a eletrodo de vidro	24
2.3 Dificuldades em medições de pH	26
2.3.1 pH médio	26
2.3.2 Curva de titulação ou curva de equalização	27
2.3.3 O efeito <i>buffering</i>	28
2.3.4 A medição de pH	29
2.4 Sensores virtuais.....	32
2.5 Aplicações de sensores virtuais	38
2.6 Desenvolvimento de sensores virtuais.....	42
3 METODOLOGIA PROPOSTA NESTE TRABALHO.....	45
3.1 Sistemas Nebulosos	46
3.2 Relevância da modelagem nebulosa.....	48
3.3 Modelos nebulosos baseados em regras	50
3.3.1 Modelo nebuloso linguístico	50
3.3.2 Modelo nebuloso tipo Takagi-Sugeno (TS)	51
3.4 Construção de modelos nebulosos.....	53
3.5 Clusterização ou agrupamento nebuloso	54
3.5.1 <i>Fuzzy C-Means</i> (FCM).....	57
3.5.2 Clusterização pelo algoritmo de Gustafson-Kessel.....	59
3.6 Estimação dos parâmetros consequentes das regras.....	62
3.7 Modelos nebulosos utilizados.....	63
3.7.1 Modelo de Wang-Langari.....	63
3.7.2 Modelo por “Regras Limitadas”	67
3.7.3 Modelo gerado a partir da clusterização de Gustafson-Kessel	73
3.8 Métodos de identificação de sistemas variantes no tempo.....	73
3.8.1 Mínimos quadrados recursivos com fator de esquecimento constante	74
3.8.2 Mínimos quadrados recursivos com fator de esquecimento variável	75
3.8.2.1 Algoritmo de Fortescue	75
3.8.2.2 Algoritmo de Wang-Langari	76
3.9 Equacionamento utilizado na implementação numérica	78
4 PROCESSO DE NEUTRALIZAÇÃO DE pH	79
4.1 Descrição de um processo de Neutralização de pH	79
4.1.1 Modelo com apenas uma variável de saída (pH)	81
4.1.2 Modelo com duas variáveis de saída (pH e nível).....	84
4.2 Desenvolvimento de um modelo virtual para o processo de neutralização de pH	89
4.2.1 Intervalos de operação e indicação do modelo.....	92

5 RESULTADOS E DISCUSSÃO	98
5.1 Aplicações do método proposto	98
5.2 Sensor virtual aplicado ao processo de neutralização de pH.....	108
6 CONCLUSÃO E CONSIDERAÇÕES FINAIS	114
BIBLIOGRAFIA CONSULTADA	118
APÊNDICE A – Modelo virtual do processo de neutralização de pH feito em Matlab®/Simulink®.....	127
APÊNDICE B – <i>Script-file</i> em Matlab® para aplicação ao primeiro exemplo do item 5.1	130

1 INTRODUÇÃO

Existem problemas no mundo real para os quais inferência de variáveis, detecção e classificação de falhas, previsão de desempenho, aproximação de funções não-lineares e identificação de padrões são requisitados. Nestes tipos de problemas, as soluções tornam-se caras e complexas, se modelos tradicionais são os únicos a descrever o comportamento do sistema. Isto ocorre devido a vários fatores: tipicamente há grande quantidade de informação envolvida no problema; as limitações de um modelo determinístico em generalizar, ou seja, tirar conclusões gerais dos dados apresentados; a vulnerabilidade dos modelos empíricos convencionais aos sinais de perturbação ou às novas condições do processo; a complexidade do modelo descrito; ou simplesmente por causa da não existência do modelo (ASCENCIO; GALICIA, 2000).

Uma das características de muitos processos industriais é a complexa inter-relação entre as variáveis do processo, o que tem levado ao desenvolvimento de técnicas que permitam a estimação de determinadas variáveis através de informações adquiridas por meio de medições de outras variáveis.

Apesar de medições em tempo real (*on-line*) serem bastante desejadas, há duas barreiras para a viabilidade de medições *on-line* das variáveis do processo: a ausência de métodos de detecção apropriados e o alto custo dos métodos conhecidos.

Isto, com frequência, conduz ao monitoramento destas importantes variáveis do processo através do uso de análises *off-line* em laboratório, que significam perda de densidade de informação, demora na obtenção dos resultados e normalmente requerem aumento do esforço humano, resultando no desvio do processo das condições de operação desejada, produzindo uma variabilidade indesejável e uma redução no rendimento (HERNÁNDEZ; ASCENCIO; GALICIA, 1998).

Estes efeitos adversos podem não ser superados em grau aceitável pelo uso dos algoritmos de controle avançado existentes. Esforços a fim de se aliviar este problema incluem o desenvolvimento de sensores virtuais (estimadores inferenciais). Há muitas variáveis medidas *on-line* e que são amostradas de forma relativamente frequente.

Estas variáveis estão indiretamente relacionadas às variáveis difíceis de serem medidas.

Qualquer sistema de controle ou monitoramento requer o emprego de elementos de interface com o mundo real, ou mundo físico. Assim, um processo industrial requer uma diversidade de sensores para poder observar e identificar seu estado atual e poder tomar decisões de controle. Por exemplo, em um processo de produção, é extremamente importante saber se o sistema está, realmente, produzindo nas condições que foi projetado ou especificado, por isso, a utilização de sensores para sua verificação é necessária.

Dado que as áreas de aplicação dos sensores são muito diversas, existem sensores de muitos tipos, que são especificados conforme a natureza da variável que será medida e das condições da aplicação. Nos processos industriais podem haver tipos distintos de sensores medindo diferentes variáveis que são relevantes ao desenvolvimento do processo, ou seja, ao controle ou monitoramento do processo. E como processar os dados dos sensores e fazer deduções acerca do processo com base nestas informações é motivo de muito estudo, e existem teorias de controle para se aplicar a informação segundo a natureza do processo, para o caso do processo de produção, podem ser utilizados modelos multivariáveis do processo e teorias de controle adequadas para esses modelos.

Porém, a medição em tempo real de algumas variáveis importantes ao controle de certos processos industriais pode ser impraticável. Isto pode acontecer por diversas razões, como situações nas quais o sensor adequado não existe ou é proibitivamente caro, pelo que se desejaria uma maneira de se estimar a informação desta variável. Em muitos casos, a informação sobre essa variável não mensurável existe, porém em outras formas. Por exemplo, se sabe que em processos onde se tem as variáveis pressão e temperatura, as mesmas estão relacionadas, ou seja, conhecendo ou medindo uma delas, é possível inferir a outra mediante um procedimento adequado. No caso de processos de fermentação, o monitoramento da concentração de biomassa e/ou produtos secundários é essencial; todavia, não existe um sensor para medi-las em tempo real, restando basicamente duas formas de se obter informações ou medidas destas variáveis:

a) o método automático através de um cromatógrafo ou um espectrofotômetro (densidade óptica), para utilização em aplicações desta natureza, que tem um custo da ordem de US\$ 10.000,00 (US BIOSOLUTIONS BRASIL);

b) o método manual através de amostras coletadas periodicamente e analisadas por um especialista. Esta técnica, além de não fornecer sinais contínuos, devido ao atraso inerente da análise, não os fornece também com suficiente frequência, não sendo, portanto, completamente satisfatória.

Para processos onde ocorre o exposto acima, ou seja, aqueles em que é difícil medir a variável desejada devido ao sensor disponível ser, por exemplo, muito caro, demasiado lento ou inexato para a aplicação específica, surge a seguinte questão: como efetuar esta medição?

Em muitos tipos de processos industriais, onde se pode contar com uma multiplicidade de sensores e se tem controle ou monitoramento realizado por computador, baseados na ideia de inferência de informação, surgem os chamados sensores por *software* ou *soft-sensors* ou ainda sensores virtuais, os quais consistem em um modelo que estima, em tempo real, a variável desejada a partir de dados medidos da planta, ou seja, programas responsáveis por fazer a inferência tomando como base a informação existente. Os programas podem consistir em um modelo matemático de como fazer a inferência, em um modelo heurístico ou um modelo inteligente e, portanto, na obtenção destes modelos, são usados, como dados de entrada, os valores das variáveis que influenciam a variável desejada (ASCENCIO; HERRERA, 1998).

Há duas metodologias usuais que podem ser adotadas na construção de um modelo inferencial. Primeiro, a chamada *bottom up*, uma proposição teórica pode ser tomada, que é geralmente a abordagem preferida para tarefas de modelagem em engenharia. Contudo, ela requer um entendimento das propriedades físicas e químicas básicas do processo. Na indústria química, por exemplo, um entendimento incompleto do processo frequentemente inviabiliza esta abordagem. A outra alternativa comum é abordar o problema pela metodologia *top down*, gerando um modelo de entrada-saída do processo com base nos dados coletados da planta. Esta técnica é também conhecida

como metodologia de modelagem caixa-preta ou caixa-cinza (ESMAILY-RADVAR, 2001).

A identificação de modelos nebulosos é desenvolvida baseada na teoria dos conjuntos nebulosos proposta por Zadeh (1965). O principal interesse tem sido na construção de modelos nebulosos que são expressos por um conjunto de proposições linguísticas nebulosas derivadas da experiência de operadores experientes (especialista) ou de um grupo de dados observados de entradas-saídas. Contudo, para alguns sistemas muito complexos, é quase impossível estabelecer tal modelo nebuloso, devido à grande quantidade de proposições nebulosas e a altamente complicada relação nebulosa multidimensional.

Takagi e Sugeno (1985), 4217 citações segundo o *ISI Web of Knowledge*, propuseram um novo tipo de modelo nebuloso que tem provado ser efetivo na superação de algumas dessas dificuldades. Seu modelo nebuloso consiste em implicações nebulosas cujos consequentes são descritos por funções de entrada-saída lineares exatas (*crisp*). Uma outra importância de seus modelos nebulosos, segundo Jin et al. (1995a), é que, como todo o parâmetro consequente é identificado por certos algoritmos, tais como o método dos mínimos quadrados, o modelo nebuloso estabelecido é mais sistemático e objetivo. Mas, infelizmente, o procedimento de identificação é bastante complicado e é executado *off-line* (embora Sugeno e Tanaka (1991) tenham sugerido um algoritmo de identificação sucessiva, ainda há dificuldades para a implementação em tempo real), o que o torna inaceitável para tratar sistemas variantes no tempo.

Jin et al. (1995b) propuseram então o casamento de redes neurais com a teoria de conjuntos nebulosos, devido à poderosa habilidade de mapeamento e aprendizado não-linear das redes neurais artificiais, para realizar o ajuste das funções de pertinência nebulosa e modificações das regras nebulosas, tornando-o útil para projetar modelos nebulosos adaptativos e controladores nebulosos auto-organizados.

Eles afirmam que a identificação de parâmetros não pode ser efetivamente realizada usando métodos de sistemas nebulosos convencionais e, para enfrentar esta situação, é necessário desenvolver uma ferramenta matemática mais sofisticada para sistemas nebulosos. Sugerem, então, uma possível e talvez a mais esperada forma que,

de maneira simples, seria combinar a teoria de redes neurais artificiais com a teoria de conjuntos nebulosos.

Estes sistemas híbridos dão a impressão de terem as seguintes características:

- 1) conjuntos nebulosos são usados para criar uma perspectiva de percepção relevante, que possua significado físico muito claro.
- 2) todas as regras nebulosas são expressas por um grupo de pesos de uma rede neural e podem ser ajustadas de uma forma mais efetiva.
- 3) a característica não-linear da rede neural provê o modelo nebuloso de maior habilidade para descrever sistemas complexos.

O objetivo deste trabalho é desenvolver um modelo inferencial ótimo usando técnicas de identificação, através da construção de modelos nebulosos dos dados coletados, para aplicação em um processo industrial complexo, como por exemplo, um processo de neutralização de pH. A identificação de modelos nebulosos baseados em regras usando dados medidos do processo requer a identificação das variáveis de entrada e saída, da estrutura do antecedente e do consequente, das funções de pertinência, e outros parâmetros associados com a estrutura do modelo particular.

Neste trabalho, estes dados são organizados e usados na geração dos modelos nebulosos que melhor se ajustam à informação coletada, de acordo com vários métodos. Ou seja, são gerados modelos nebulosos do tipo entrada-saída. O modelo é em tempo discreto e pode ser representado por uma função F :

$$y_f(k) = F(u_1(k-1), \dots, u_1(k-N_1), \dots, u_p(k-1), \dots, u_p(k-N_p), y_f(k-1), \dots, y_f(k-N_f)) \quad (1)$$

onde $y_f(k)$ representa a saída no instante k , $u_p(k-N_p)$ são as N_p regressões da p -ésima entrada e $y_f(k-N_f)$ são as N_f regressões da saída, ou seja, as estimativas de instantes anteriores são usadas de forma a representar a dinâmica do sistema (OLIVEIRA; GARCIA, 2003).

Nessa proposta, as relações entre entrada e saída presentes e futuras do processo são definidas a partir de sentenças do tipo SE-ENTÃO, de forma que se possa entender

ou mesmo prever o comportamento dos sistemas a partir de suas entradas e saídas. Este método de modelagem está intimamente relacionado ao domínio da Inteligência Artificial e, como tal, às vezes é considerado um modo especial de modelagem que resulta de herança do comportamento humano, ou seja, como um método para imitar a inteligência humana sendo capaz de resolver o problema onde o experimentador tem apenas um conhecimento intuitivo do processo.

Na prática modelos usando lógica nebulosa são meras equações matemáticas que mapeiam entradas $\mathbf{u}$ em saídas $\mathbf{y}$, isto é, eles formam um mapeamento multidimensional de entradas $\mathbf{u} \in \mathbb{R}^n$ para saídas $\mathbf{y} \in \mathbb{R}^m$ via uma função não-linear $f_{NL} : \mathbb{R}^n \rightarrow \mathbb{R}^m$ com parâmetros θ ou

$$\mathbf{y} = f_{NL}(\mathbf{u}, \theta) \quad (2)$$

O funcional f_{NL} pode ser um filtro NFIR (*nonlinear finite impulse response*) estático, ou um mapeamento NARX (*nonlinear autoregressive exogenous*) ou NARMAX (*nonlinear autoregressive moving average with exogenous inputs*) ou um sistema dinâmico com saídas regressivas (VAN GORP; SHOUKENS, 1999).

Como o enfoque principal deste trabalho é poder modelar sistemas variantes no tempo, um cuidado especial deve ser tomado, levando em consideração modelos baseados em lógica nebulosa, quando da identificação dos parâmetros dos consequentes, que, além das dificuldades já existentes na estimação dos parâmetros, ainda apresenta o problema com a variação dos mesmos em tempo real, o que requer a utilização de técnicas de identificação que possam ser usadas em sistemas variantes no tempo. Em particular, neste trabalho, deseja-se construir um sensor virtual para estimar a medida do pH do processo CSTR (*Continuous Stirred Tank Reactor* - Reator de tanque agitado de fluxo contínuo) levando em consideração sua não-linearidade e sua variância no tempo. Este processo possui diversas aplicações científicas e industriais. Apesar da medida do pH, em tempo real, poder ser obtida através de sensores comuns no mercado, este processo é usado aqui por ser altamente não-linear e variante no tempo. Além disso, foi desenvolvido um modelo virtual para este processo e implementado usando Matlab/Simulink, baseado nos estudos realizados na

Universidade de Santa Bárbara, Califórnia, onde foi desenvolvido um experimento em escala de laboratório (HALL; SEBORG, 1989) (WALLER; MÄKILÄ, 1981). Para lidar com este tipo de processo, desenvolveu-se também uma planta piloto na Universidade Nacional de San Juan, Argentina (ALVAREZ et al., 2001).

Este trabalho está dividido em 6 capítulos como descritos a seguir.

No capítulo 2, será feita uma descrição de sensores de pH, suas aplicações e seus problemas e, também, é dada uma visão geral das aplicações de sensores virtuais encontrados na literatura, na indústria de processos e na indústria de instrumentação.

No capítulo 3, será descrito o método utilizado para construção de um sensor virtual utilizando modelos nebulosos, fazendo um resumo sobre sistemas nebulosos, a relevância da modelagem nebulosa, tipos de modelos nebulosos, e sua construção.

O capítulo 4 apresentará os resultados obtidos no desenvolvimento de um modelo virtual, usando o *software* Matlab®/Simulink®, baseado em um modelo de referência *benchmark model* desenvolvido pela Universidade de San Juan – Argentina, a partir de estudos do processo de neutralização de pH realizado pela Universidade de Santa Barbara – USA, modelo virtual que será usado para testar o desempenho do sensor virtual.

No capítulo 5, será feita uma análise dos resultados obtidos na utilização do método proposto, inicialmente em identificação de sistemas complexos, empregando-o em exemplos encontrados na literatura de identificação de sistemas e, também, na aplicação como sensor virtual na planta virtual de neutralização de pH.

O capítulo 6 apresenta as conclusões e considerações finais deste trabalho.

2 SENSORES VIRTUAIS DE pH

2.1 Definição de pH

O pH muitas vezes é considerado como uma quantidade com a mesma natureza que massa e energia, com as quais balanços de conservação podem ser estabelecidos. Esta concepção, contudo, pode conduzir a sérios erros.

Durante sua pesquisa visando melhorar a qualidade da cerveja, o bioquímico dinamarquês, P. L. Sorensen, criou, no início do século 20, o conceito de “potencial hidrogeniônico” – pH (FELTRE, 2001). O pH é uma medida da concentração dos íons de hidrogênio, $[H^+]$, em uma solução aquosa e serve para indicar sua acidez (ou alcalinidade). Tendo em vista que o número de mols¹ de H^+ por litro da solução é um número ínfimo, expresso em potências negativas de 10, Sorensen optou, por razões práticas, por expressá-la apenas pelo expoente de 10 (e sem o sinal negativo). Assim, por exemplo, é de 5,3 o pH de uma solução com $10^{-5,3}$ moles de H^+ por litro. Matematicamente, o pH é calculado como segue:

$$pH = -\log_{10}[H^+] \quad (3)$$

A água se dissocia em íons H^+ e íons OH^- (hidroxila), em qualquer temperatura, segundo a reação de equilíbrio a seguir:



A concentração molar de íons H^+ na água destilada², à temperatura de 25° C, é de $10^{-7,0}$ e serve como referência fundamental para o pH; ou seja, 7,0 é o valor do pH neutro. Concentrações crescentes acima desta referência indicam acidez crescente, ao passo que concentrações decrescentes abaixo desta referência indicam alcalinidade crescente (FELTRE, 2001). A Figura 1 mostra a escala do pH (absoluta e relativa) ilustrada com diversos produtos do dia-a-dia. Embora raras e extremamente reagentes,

¹ Um mol é um número especial e vale $6,022 \times 10^{23}$. A unidade Molar, M, medida de concentração, é, por definição, o número de mols/litro (<http://www.chemistrycoach.com>).

² Água destilada significa água livre de metais pesados, tais como lítio, magnésio, etc.

existem ácidos e bases fortíssimos, com pH negativo e superior a 14, respectivamente (LUCENA, 2006).

		Concentração relativa de H^+	Exemplos
	0	10.000.000	Ácido de bateria, ácido fluorídrico forte, HCl 1 M
	1	1.000.000	Ácido de clorídrico, ácido segregado pela mucosa estomacal
	2	100.000	Suco de limão, vinagre, suco gástrico
	3	10.000	Suco de laranja, suco de uva
	4	1.000	Suco de tomate, chuva ácida
	5	100	Café preto, água potável
	6	10	Saliva, urina
Acidez ↑	7	1	Água destilada
Neutro	8	1/10	Água do mar
	9	1/100	Clorox (produto de limpeza)
	10	1/1.000	Leite de magnésia
	11	1/10.000	Solução de amônia (NH_3)
Alcalinidade ↓	12	1/100.000	Água de sabão
	13	1/1.000.000	Limpa forno
	14	1/10.000.000	NaOH 1 M (1 mol/litro) (Soda cáustica)

Figura 1 - Ilustração da escala do pH e alguns exemplos com substâncias comuns.

Normas de saneamento determinam que o pH da água fornecida à população pelas companhias de abastecimento tem de encontrar-se na faixa de 7,2 a 7,6 (BERNARDO, 2005). As empresas responsáveis pelo tratamento e abastecimento da água em centenas de cidades brasileiras, até onde se sabe, realizam manualmente o controle deste pH, que pode alterar-se nos reservatórios, em decorrência de fatores tão diversos como chuva ácida, arraste por águas pluviais de material orgânico em decomposição, poluição por atividade humana, dentre outros. Este controle manual pode levar a custo operacional maior, além de menor garantia da qualidade da água fornecida à população.

Por outro lado, o pH não é uma medida total de acidez ou alcalinidade de uma solução, mas somente seu grau de ionização. Assim, bases fortes e ácidos fortes se ionizam quase completamente quando dissolvidos em água. Dessa forma, se quantidades equivalentes de ácido forte e base forte fossem adicionadas em água, o resultado seria uma solução neutra, onde as hidroxilas da base iriam se combinar com

os hidrogênios do ácido. Por outro lado reagentes fracos somente ionizam-se parcialmente. Se, por exemplo, uma quantidade equivalente de ácido forte e base fraca, ou ácido fraco e base forte, forem misturados, a solução final não será neutra. Como resultado desse fenômeno, mais quantidade do reagente forte será requerido para tornar a solução neutra. Portanto, as palavras “forte” e “fraco” neste contexto estão referidas à concentração ou ao grau de pureza do material (ácido ou base).

Portanto, são necessárias mais informações sobre a química do sistema, ou seja, sobre as reações químicas na mistura. Concluindo, o pH de uma mistura não pode ser calculado se só são conhecidos os valores de pH (e as proporções de mistura) das soluções individuais (SOTOMAYOR, 1997).

2.2 Sensor de pH a eletrodo de vidro

Estima-se que atualmente existam milhões de pH-metros operando no mundo, em aplicações que vão das indústrias química, petroquímica, agro-alimentar e farmacêutica, a biologia e clínica. Apesar do desenvolvimento recente de novos sensores de pH (SAFAVI, BAGHERI, 2003; SHARMA, GUPTA, 2003), sua medição continua sendo realizada, em esmagadora maioria, por meio de sensor a eletrodos de vidro, descrito pela primeira vez por Haber, em 1909 (ASCH, 1999). A Figura 2 mostra o diagrama esquemático de uma célula eletroquímica para medição de pH. Como se vê, há dois eletrodos: um chamado de eletrodo de medição e outro, de referência. Ambos os eletrodos são preenchidos com uma solução de pH conhecido (geralmente, com cloreto de potássio, KCl, com pH igual a 7,0).

O eletrodo de medição é formado por um tubo de vidro impermeável, de alta resistência elétrica, ao qual se solda em sua extremidade inferior uma membrana de vidro dopado com íons de lítio (COBBOLD, 1974; ALLABOUTCIRCUITS, 2010), de formato geralmente esférica, cilíndrica ou cônica. Esta membrana é semipermeável, ou seja, neste caso, permeável apenas ao íon H^+ . Assim, dado que dentro do tubo de vidro há uma solução com pH conhecido (diga-se, com concentração molar de íons H^+ conhecida), caso a concentração molar de íons H^+ da solução que preenche o recipiente externo seja diferente daquela, acontecerá movimentação de íons H^+ do lado

de maior para o de menor concentração. Esta migração iônica resultará, no equilíbrio dinâmico, em uma diferença de potencial através da membrana de vidro proporcional à diferença entre os pHs das soluções. Quantitativamente, calcula-se esta diferença de potencial por meio da equação de Nernst (ASCH, 1999; COBBOLD, 1974; NATIONAL INSTRUMENTS, 2005)

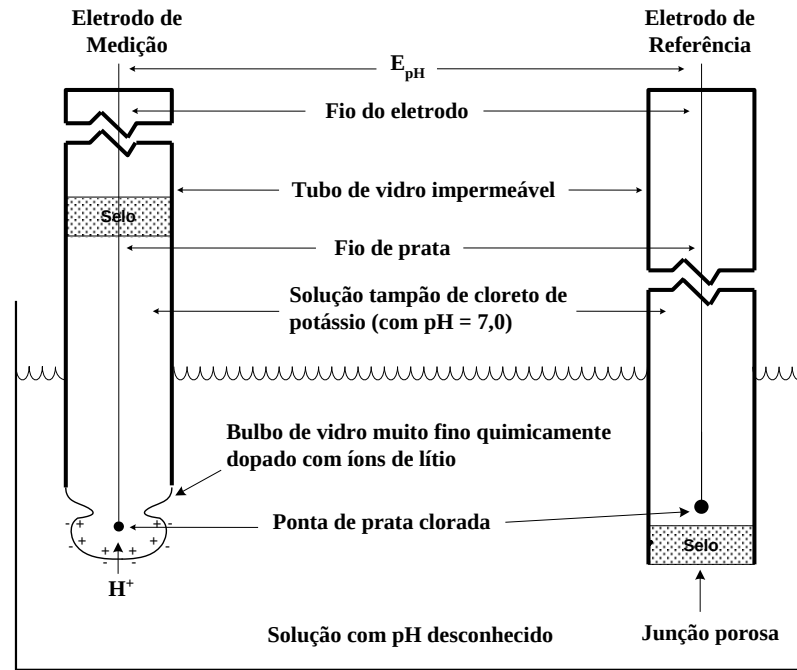


Figura 2 - Ilustração dos eletrodos de vidro para medição de pH (LUCENA, 2006).

$$E_{pH} = E_0 + \frac{2,3RT}{nF} \cdot \log[H^+] \quad (5)$$

Onde

E_{pH} é a tensão através da membrana de vidro;

E_0 é a diferença de tensão em uma solução de pH igual a 7,0;

R é a constante universal dos gases ($8,3143 \text{ JK}^{-1}\text{mol}^{-1}$);

T é a temperatura em Kelvin;

n é o número de valência dos elétrons por mol (igual a 1, para o H^+);

e F é a constante de Faraday (96487 Cmol^{-1}).

2.3 Dificuldades em medições de pH

2.3.1 pH médio

Como é possível observar através dos relatos dos itens anteriores o conceito de pH não é trivial. Por exemplo, muitos trabalhos têm-se referido a métodos para calcular “valores de pH médio”, provocando muitas discussões e confusões. Parte da confusão tem sido causada pelo fato de que o conceito não tem sido bem definido.

O “pH médio” é descrito como o pH resultante de uma mistura de soluções. Pela definição apresentada anteriormente, pela Equação 3, pode-se pensar que o “pH médio” de algumas soluções pode ser obtido pela média do pH obtido depois da mistura das soluções, ou seja, da seguinte maneira:

$$p[H]_{AV} = \frac{1}{n} \sum_{i=1}^n \{pH\}_i \quad (6)$$

Gustafsson e Waller (GUSTAFSSON; WALLER, 1986) descrevem sobre um outro método sugerido na literatura que, para uma mistura de iguais volumes (ou fluxo), o pH médio é calculado por:

$$p[H]_{AV} = -\log_{10} \left\{ \frac{1}{n} \sum_{i=1}^n [H^+]_i \right\} \quad (7)$$

Esta maneira é sugerida como uma correção do procedimento descrito pela Equação 6. Ambos os métodos de cálculo, Equações 6 e 7, são incorretos. Por exemplo, para uma mistura de quatro soluções de igual volume, com valores de pH de 2, 4, 6 e 8 respectivamente, temos que, aplicando a Equação 6, o $p[H]_{AV} = 5$, e aplicando a Equação 7, o $p[H]_{AV} = 2,6$.

Nenhum dos resultados é correto. Se um ácido forte com pH=4 é diluído em água pura, o pH final não é independente da presença de um ácido fraco (*buffer*) no sistema. Portanto, são necessárias mais informações, especificamente, sobre a química do

sistema, ou seja, sobre as reações químicas na mistura. Em conclusão, o pH de uma mistura não pode ser calculado se só são conhecidos os valores de pH (e as proporções de mistura) das soluções individuais (GUSTAFSSON; WALLER, 1986).

2.3.2 Curva de titulação ou curva de equalização

A unidade primária de informação básica requerida para o projeto de sistemas de controle de pH é a curva de titulação, algumas vezes referida como a curva característica ou a curva de equalização. A forma da curva é conhecida da relação logarítmica mostrada na Equação 3. A identificação de cada espécie no fluxo de saída, suas concentrações, suas constantes de ionização (ou equilíbrio), permite calcular a curva de titulação. A Figura 3 mostra as curvas de titulação para diversos processos ácidos-base.

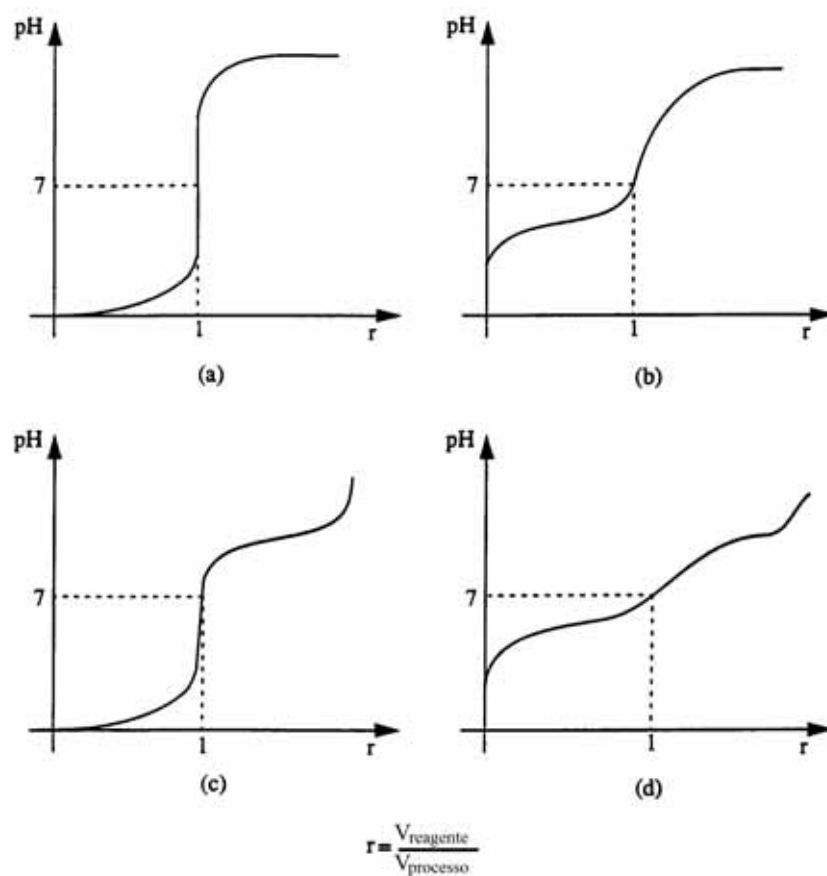


Figura 3 – Curvas de titulação para: a) ácido forte – base forte, b) ácido fraco – base forte, c) ácido forte – base fraca, d) ácido fraco – base fraca.

Ponto neutro – na curva de titulação é o ponto no qual a concentração de íons de hidrogênio é igual à concentração de íons hidroxila.

Ponto de equivalência – na curva de titulação é o ponto no qual a concentração de íons ácido é igual à concentração de íons base.

O ponto de equivalência coincide com o ponto neutro para um sistema ácido forte – base forte. A inclinação da curva de titulação é infinita no ponto de equivalência.

As palavras, forte e fraco dizem respeito à concentração ou grau de pureza do material (ácido ou base). Por exemplo, uma solução ácida de 95% é mais forte que uma solução ácida de 4%. Na Figura 3, esta definição é usada para mover a curva de titulação horizontalmente para direita se o ácido torna-se forte (em concentração), ou para a esquerda se o ácido torna-se fraco. Este deslocamento da curva, no entanto, provoca uma severa complexidade no controle de pH (GUSTAFSSON; WALLER, 1986) (SOTOMAYOR, 1997).

2.3.3 O efeito *buffering*

Se quantidades equivalentes de uma base forte e um ácido fraco, ou um ácido forte e uma base fraca, são misturados em um tanque, a solução resultante pode não ser neutra. Como resultado deste fenômeno, será necessário uma maior quantidade de reagente forte para neutralizar o reagente fraco. Por exemplo, uma amostra de água residual com $\text{pH}=2$ irá requerer uma maior quantidade de reagente básico para ser neutralizada se a acidez é causada por ácido acético (um ácido fraco) que se ela fosse causada por ácido clorídrico (um ácido forte). Este efeito é chamado *buffering*, que é a capacidade de uma solução para resistir a mudanças no pH.

Reagentes fracos mantêm uma reserva de ácido ou base não dissociados. Considere-se um ácido fraco sendo neutralizado por uma base forte. Nos valores baixos de pH, o ácido tem uma grande capacidade de absorve íons hidroxila da base devido à reserva de ácido não dissociado. Como a base é titulada na solução, íons hidrogênio são removidos dela. Porém, como os íons hidrogênio são neutralizados, o

equilíbrio da reação varia, liberando mais íons hidrogênio e provocando uma pequena mudança no pH, desta maneira, as soluções *buffered* resistem às mudanças no pH.

O uso de um elemento *buffer* em um processo ácido forte – base forte, como o da Figura 4, provoca o efeito *buffering* na curva de titulação, o que significa que a curva é menos sensível perto do ponto de neutralidade, e o baixo ganho do processo facilita o controle de pH. Contudo, além do fato de facilitar o controle, que é bastante desejado em processos de neutralização de pH, o *buffering* torna o sistema variante no tempo (GUSTAFSSON; WALLER, 1986) (HALL; SEBORG, 1989) (HENSON; SEBORG, 1994) (SEBORG; EDGAR; MELLICHAMP, 2004).

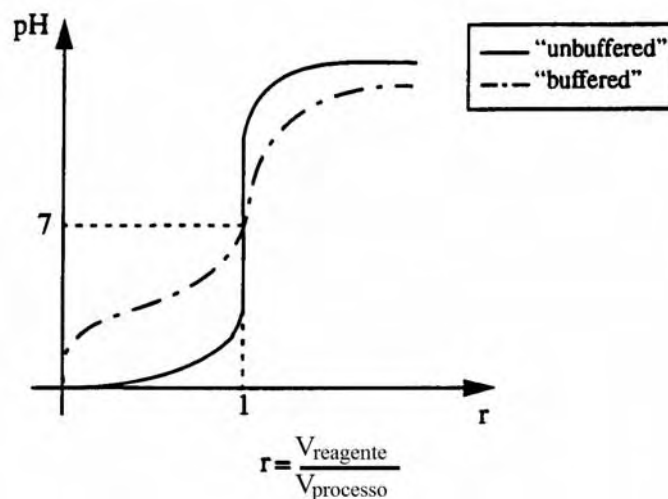


Figura 4 – Curvas de titulação para soluções *buffered* e *unbuffered*.

2.3.4 A medição de pH

É comum, em indústrias que possuem sistemas de medição de pH, ocorrerem situações do tipo:

“Algumas vezes por semana, o pessoal de instrumentação recebe um chamado para a correção em uma falha de alguns dos sistemas de medição de pH da planta; este logo pega o seu ‘kit de sobrevivência’ que contém minimamente os *buffers* de calibração, soluções químicas de limpeza, água em uma pisseta de laboratório e papel higiênico e segue para o processo, todo equipado e munido de uma grande boa vontade para dedicar alguns minutos ou horas à resolução do problema. Após desmontar o eletrodo

de seu suporte, ele testa o eletrodo, tenta todo tipo de limpeza química e, muitas vezes, acaba descobrindo que o problema vem de uma baixa isolamento do cabo, que foi atacado por corrosão ou umidade, levando então à troca de cabo e eletrodo, pois o ponto de medição está fora do ar já faz algum tempo e a sala de controle já está impaciente com a falta da medição de pH por tanto tempo no supervisório (SABADIN, 2008).

É por este motivo que sistemas de pH são a grande dor de cabeça das equipes de instrumentação. A raiz dos problemas é que tudo funciona perfeitamente até que ao primeiro sinal de perda de isolamento do cabo entre o sensor e transmissor, causada por diversos fenômenos como umidade, sujeira, oxidação de terminais e muitos outros, ocorre uma variação da resistência deste mesmo cabo, que a partir deste momento entrega ao transmissor um valor em milivolts com um erro, erro este que acarreta um desvio do valor calculado de pH, mas que não é uma variação real de pH do processo, mas sim um simples erro de medição. A partir deste momento, é iniciada uma reação em cadeia, onde os sistemas automáticos de correção de pH começam a atuar, abrindo e fechando válvulas e trabalhando de forma a corrigir uma variação de pH que não é real e só leva a erros maiores, que serão percebidos somente na próxima intervenção preventiva, ou no momento em que o erro se torne tão grande que alguém de processo perceba e intervenha corretivamente (McMILLAN, 1994).

Neste momento o instrumentista é chamado para compensar o erro, trocando o cabo e/ou limpando terminais, se este for o caso crítico de desvio. Se for um caso mais simples, ele promove uma nova calibração, para que os efeitos desta variação sejam compensados no conjunto transmissor, cabo e eletrodo, de forma que o erro desapareça momentaneamente, até que o próximo fenômeno altere novamente a condição do cabo e novos erros apareçam. É importante notar que a calibração foi uma ação corretiva eficaz, pois compensou o erro, mas foi motivada por algo externo ao eletrodo e seus componentes internos. A percepção é que um dos grandes vilões na medição de pH não é o eletrodo em si, e sim os problemas que aparecem no caminho entre o sensor e o transmissor de pH, ou seja, a forma como este sinal (mV) delicado trafega até o transmissor.

Pensando nisto, os fabricantes de sistemas de medição desenvolvem conectores com banho de ouro, blindagens especiais dos condutores e complexas proteções contra umidade, com o intuito de minimizar o impacto devastador deste efeito; porém, isso resulta em cabos que devem ser extremamente bem cuidados e protegidos em eletrodutos individuais, com poucos metros de comprimento e exigindo calibrações freqüentes de todo o conjunto no local (SABADIN, 2008). A Figura 5 mostra a configuração encontrada em um sistema analógico de medição de pH.



Figura 5 – Cadeia de medição analógica (SABADIN, 2008).

O conceito do sensor digital é baseado na idéia de fazer todos os processamentos críticos ainda dentro do sensor, converter os dados de analógico para digital ainda dentro do mesmo e só então estabelecer uma conexão entre sensor e transmissor, livre de interferências e livre de conectores metálicos passivos de efeito da umidade e corrosão. A Endress+Hauser, uma das empresas líderes em instrumentação e soluções em automação para processos industriais, desenvolveu um sistema digital de medição de pH, com uma pequena eletrônica no próprio sensor com quase tudo o que existia num transmissor analógico, com o propósito de conseguir um diagnóstico ativo, o fim da conexão de alta impedância, e menor quantidade de intervenções. A Figura 6 mostra a cadeia de medição do sistema digital.

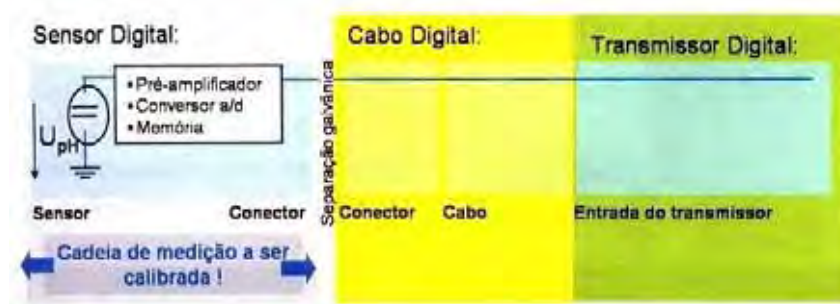


Figura 6 – Cadeia de medição digital (SABADIN, 2008).

Neste novo conceito, alguns sensores poderão estar disponíveis no laboratório já limpos, tratados e pré-calibrados, ou seja, prontos para substituir um determinado sensor que esteja terminando seu ciclo de operação. Assim, todas as soluções de calibração, limpeza e demais ferramentas de uso comum do especialista de instrumentação analítica ficam dentro do laboratório, resumindo o trabalho em campo a uma troca do sensor. Isto torna o sistema de medição muito melhor e mais simples, mas mesmo com todos esses avanços o processo fica sem o sensor no momento de sua troca ou manutenção.

2.4 Sensores virtuais

Desde 1999, anualmente é realizado um seminário internacional, patrocinado pelo IEEE (*Institute of Electrical and Electronics Engineers*), em Sistemas de Medições Virtuais e Inteligentes (VIMS), no qual apresentam-se trabalhos em tecnologias emergentes em instrumentação, medição e em sistemas virtuais para instrumentação e medição. Como relatado em um artigo apresentado na VIMS/2001 por Siegel (2001), apesar de ambos os tópicos Instrumentos Virtuais e Sensores Virtuais serem temas destes seminários, este último historicamente quase não é descrito nos programas, apesar de sua relevância e aplicação industrial ser observada.

Em trabalhos de Dorst (1998), são encontrados alguns exemplos do termo “sensor virtual”, que também é chamado de “sensor lógico” introduzido por Henderson e Shilcrat (HENDERSON, 1984 apud DEKHIL, 1998).

A noção de “sensor lógico” é descrita como uma metodologia para especificar qualquer sensor de tal forma que esconda sua natureza física, tendo como principal objetivo desenvolver uma apresentação coerente e eficiente da informação fornecida por muitos sensores de diferentes tipos.

Um sensor lógico não tem que ser uma entidade física; pode ser um programa, ou uma combinação de *software* e *hardware* e, em qualquer caso, deve ser visto como um sensor distinto que poderia ser substituído por um *hardware* especial. Ou seja, uma combinação (real ou imaginária) de um sensor físico com processamento e

interpretação de dados por *software*, capaz de medir um dos relevantes parâmetros que caracterizam um sistema.

A seguir apresenta-se uma descrição bastante formal dada por Groen (DORST, 1998) e transcrita por Siegel (SIEGEL, 2001) sobre o “sensor virtual”:

Em qualquer sistema cujo objetivo seja o sensoriamento, um *sensor virtual* é um dispositivo (conceitual) cuja saída pode ser modelada em termos dos parâmetros relevantes e das saídas de outros sensores reais. Os módulos de sensores virtuais devem ser escolhidos no mais alto nível de abstração que permita uma caracterização suficientemente exata do comportamento geral do sistema, mas no qual as interações entre os vários módulos de sensores virtuais sejam (relativamente) simples, tanto em suas (in)dependências estatísticas quanto em suas relações causais. Em um sistema simulado, há a exigência adicional de que os modelos dos sensores virtuais devam ser passíveis de validação.

Siegel (2001) apresenta também outra definição bastante abrangente que é coerente com as utilizadas em diversos trabalhos:

O termo ‘sensor virtual’ é realmente usado de duas maneiras: primeiro para que possa se incluir no modelo do sistema um agente que fielmente represente um sensor real a fim de realizar uma análise de custo-benefício antes de realmente comprá-lo e instalá-lo, e segundo, para que possa se incluir no modelo do sistema um agente que represente um sensor desejado.

Assim, pode-se dizer que um ‘sensor virtual’ é um sistema projetado para substituir a indisponibilidade momentânea ou permanente de um sensor em uma planta. Esta substituição pode ser exigida porque o sensor real tenha falhado ou porque tenha sido removido para manutenção. Pode também acontecer que o sensor real seja usado de forma compartilhada (*time-sharing*), ou não tenha no mercado um sensor razoável disponível.

A Figura 7 mostra o diagrama de uma planta na qual a variável y , medida usando um sensor, está para ser substituída pelo sinal y_m fornecido pelo sensor virtual. Pode também acontecer que um SV (sensor virtual) necessite ser projetado porque não existe um sensor disponível. Em geral, o SV é o modelo de uma parte da planta convenientemente escolhida. As entradas para este modelo são outras variáveis medidas – controles ($u_1, \dots, u_r$), saídas da planta (η_1, η_2, η_3) e distúrbios medidos ($p_1, \dots, p_n$). O modelo usado pode ser do tipo entrada-saída (ARMAX) ou na forma de equações de estado e de saída. Este último caso comumente conduz ao uso de um Filtro de Kalman Estendido. Em certos casos, pode ser apropriado usar modelos

nebulosos ou modelos baseados em redes neurais, principalmente quando não se podem utilizar informações a respeito do processo, tendo-se como informações apenas os dados medidos da entrada e saída do processo, e procura-se utilizar modelos locais ao invés de globais, com o intuito de se obter modelos mais facilmente interpretáveis.

Com referência ainda à Figura 7, quando o sensor real está ativo, a chave está na posição *R*. Então o modelo do SV é identificado, enquanto a malha de controle usa o sinal y dado pelo sensor real. Quando o sinal do sensor não está mais disponível, a chave deve ser virada para a posição *V*, em cujo caso a malha de controle usa o sinal y_m fornecido pelo SV, e a identificação do modelo ou a estimação paramétrica é interrompida.

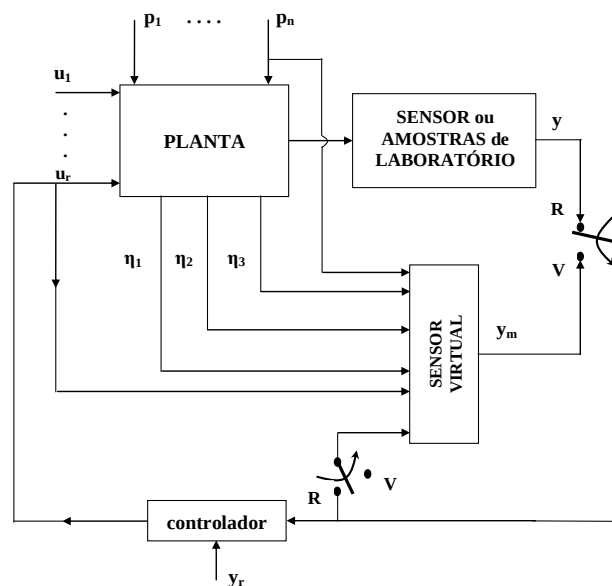


Figura 7 – Sensor virtual conectado a uma planta

Um outro importante uso potencial para os SV é quando da detecção de falha ou grave perda de calibração no sensor real, ou seja, se o sensor real estiver fora de serviço por qualquer motivo, a variável deixa de ser medida, e isso pode resultar em sérios problemas de operação da planta ou processo. Para evitar esses problemas, dependendo das condições técnicas e econômicas, pode ser utilizado como suporte um sensor real ou um SV, que geralmente é implementado por *software*.

A Figura 8 ilustra o compartilhamento no tempo de um sensor entre diferentes partes de uma planta. Para aquele ponto que é momentaneamente desconectado do

sensor real, o correspondente sensor virtual fornece um sinal estimado. Durante o intervalo de tempo em que o sensor real está efetuando a medida, os correspondentes parâmetros do modelo do sensor virtual são estimados. Esta abordagem tem sido usada, por exemplo, para estimar a distribuição do tamanho de partículas em uma planta de moagem (trituração) no Chile (GONZÁLEZ et al., 1994).

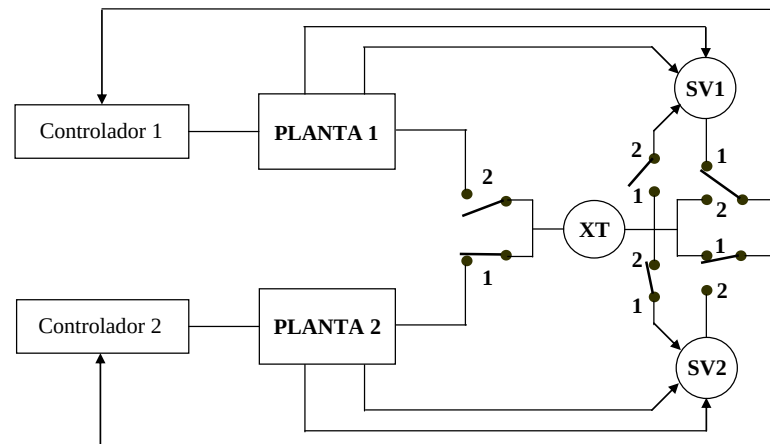


Figura 8 – Os sensores virtuais SV1 e SV2 são usados em duas malhas de controle quando o sensor XT é compartilhado no tempo entre as duas plantas.

Se um sensor não está disponível, em alguns casos, um sensor virtual pode ser desenvolvido e calibrado usando amostragem e análise *off-line* de laboratório, como ilustrado na Figura 9.

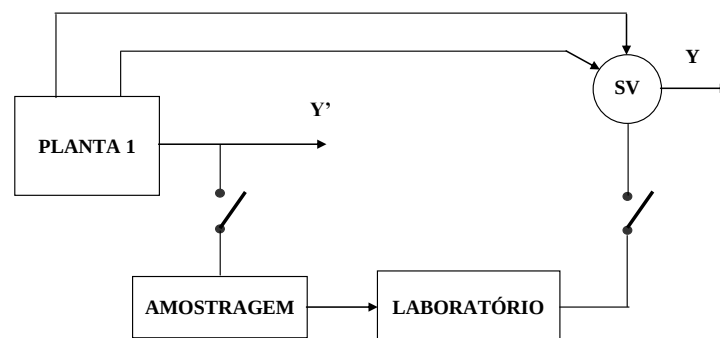


Figura 9 - Sensor virtual substituindo um sensor real inexistente

A Figura 10 mostra o caso no qual o sensor virtual usa um modelo de estados da planta, para estimar uma variável para a qual nenhum sensor está disponível. Este

procedimento tem sido usado para estimar água e minério contidos em uma retífica semiautógena (GONZÁLEZ et al., 1994).

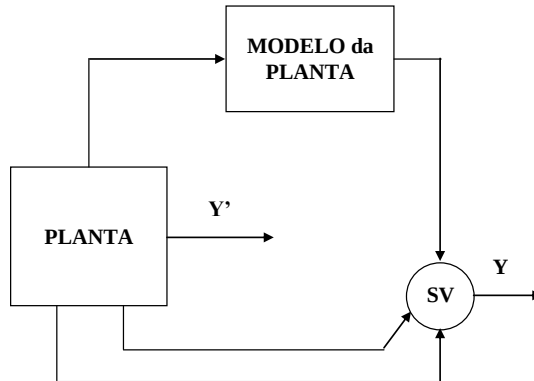


Figura 10 - Sensor virtual usando um modelo de planta quando nenhum sensor real está disponível.

Vários problemas, além daqueles de praxe em identificação de sistemas, como testes dinâmicos e coleta de dados, escolha da representação matemática a ser usada, determinação da estrutura, dentre outros, surgem no desenvolvimento do sensor virtual:

- (i) seleção de um modelo;
- (ii) escolha dos parâmetros do modelo;
- (iii) estimação da confiabilidade como uma função do tempo decorrido desde que o mesmo foi colocado em operação;
- (iv) problemas causados na malha de controle decorrentes da substituição do sensor real pelo sensor virtual.

Em geral, um modelo satisfatório para controle, baseado em modelo, não é necessariamente adequado para um sensor virtual. De fato, o modelo de controle é normalmente usado para prever um sinal dentro de um período relativamente curto, enquanto no caso de um sensor virtual, um longo período de predição é normalmente necessário. Neste contexto, por exemplo, um distúrbio de variação lenta - com respeito ao tempo de resposta da planta - pode ser considerado constante, no caso de um modelo para controle, mas pode necessitar ser representado em um modo dinâmico em muito maior extensão de tempo de predição, para o modelo de sensor virtual.

Outro aspecto a ser considerado no modelo do sensor virtual está relacionado a sua estrutura relativa à porção da planta usada para representar o sinal desejado cuja medida vai substituir. Quanto maior a parte da planta a ser modelada, maior a complexidade do modelo resultante tende a ser, e o tempo exigido para atualizar os parâmetros do modelo pode ser muito mais longo. Também um número maior de distúrbios não mensuráveis pode estar incluído, causando o aumento dos erros de estimação do sinal. Um modelo usando somente uma pequena parte da planta seria então desejável, usando medidas que estão de algum modo envolvidas com o sinal a ser predito. Mas, por outro lado, como será discutido abaixo, até mesmo se um bom modelo é obtido usando uma porção convenientemente pequena da planta, problemas sérios podem surgir quando o sensor virtual é chamado para substituir o sensor real em uma malha de controle.

Outro problema relacionado à estrutura do modelo, que pode levar um modelo de sensor virtual a ser diferente de um modelo usado em controle, está relacionado à determinação da estrutura do modelo (por exemplo, achar quais são os componentes importantes - sinais medidos, incluindo atrasos diferentes e formas funcionais diferentes) que constitui as entradas do modelo. Estas podem ser determinadas usando-se vários métodos. Em algumas plantas, onde perturbações medidas têm uma grande influência na variável a ser predita pelo modelo, pode acontecer que o controle seja eliminado da estrutura. Isto é correto para um modelo de sensor virtual, mas certamente não, para um modelo de controle, onde o componente de controle deve ser forçosamente incluído (contudo o problema de controle pode se tornar difícil devido a esta situação).

A capacidade de um sensor virtual funcionar conforme esperado, eficientemente e sem falhas, ou seja, confiavelmente, é uma outra questão a ser considerada:

- (i) no momento em que o sensor virtual está para ser colocado em operação, e
- (ii) durante o período em que ele está em operação.

No primeiro caso, a fim de descobrir se o sensor virtual está apropriado, quando ele estiver para ser colocado em operação, seu desempenho deve ser avaliado durante um período de tempo t_c antes da falha. Isto é feito, na verdade, medindo-se o erro quadrático médio entre a saída do sensor virtual e o sinal real (que está na verdade

sendo medido) durante o tempo t_c . Visto que t_c não é conhecido com antecedência, permanentemente períodos de testes sobrepostos são necessários.

No segundo caso, esta confiabilidade pode ser expressa em termos da probabilidade de o erro quadrático médio entre a saída do sensor virtual e a, agora não mensurável, variável real ultrapassar um certo limite, em um dado instante do período de operação do sensor virtual.

Técnicas de modelagem específicas para sensores virtuais usando redes neurais, conjuntos e agrupamentos nebulosos, constituem um interessante desenvolvimento a ser pesquisado.

Embora testes experimentais tenham sido efetuados em plantas industriais, muito mais pesquisas são requeridas, para garantir aplicações práticas bem sucedidas. Entre outros temas, a oportuna descoberta de um sensor falho é de suma importância:

- (i) para que os parâmetros do modelo não venham a ser incorretamente estimados, e
- (ii) a fim de trocar o sensor pelo sensor virtual a tempo.

Em geral, então, a área de sensores virtuais é um rico campo para pesquisa, que pode conduzir a aplicações industriais, pois ela fornece uma maneira de aumentar a robustez das malhas de controle, ou pelo menos proporcionar um meio mais suave de diminuição de sua qualidade, em caso de falha ou remoção do sensor (GONZALEZ et al., 2004) (FORTUNA et al., 2007).

2.5 Aplicações de sensores virtuais

Em um trabalho realizado por Muir (1990), é utilizado o termo “sensor virtual” para descrever uma etapa intermediária de fusão e análise de dados que poderia ser considerado como um substituto para um sensor hipotético. O sensor virtual desenvolvido é baseado no método de estimação de variância mínima. Ele aplicou esta abordagem, teórica e experimentalmente, ao problema de identificação de um robô cinemático, descrevendo uma abordagem que combina e processa a saída de vários sensores indiretos, mas reais, para sintetizar as assim chamadas saídas, erros e ruídos característicos de outro conjunto de sensores diretos indisponíveis, mas conceitualmente desejáveis. Os resultados sugerem que o método de estimação de

parâmetros do sensor virtual combinado com o sensor de posição visual pode ser vantajosamente aplicado para realizar a identificação cinemática de robôs.

Groen tratou de várias áreas de aplicação, dentre as quais um modelo baseado em visão robótica e outro utilizado em sistema de débito automático para cobrança eletrônica de pedágio utilizado em uma auto-estrada da Holanda (SIEGEL, 2001).

Atkinson et al. (1998) descreve um sistema de monitoramento de emissão automotivo em tempo real empregando um sensor virtual baseado em rede neural. Um sistema de rede neural foi capaz de prever a potência de saída do motor em tempo real, o consumo de combustível e as emissões, usando medições dos parâmetros do motor através dos altamente transitórios ciclos de operação do motor.

Dixon et al. (1999) relata um ambiente virtual para modelar um sistema de robô móvel múltiplo que inclui a capacidade para combinar sensores reais e virtuais. Criaram uma plataforma de software, conhecida como RAVE (*Real and Virtual Environment*), que fornece a infra-estrutura necessária para explorar o comportamento de colaboração em sistemas de múltiplos robôs móveis heterogêneos. A principal característica desta estrutura é a capacidade do mesmo programa executar em um robô real ou um robô simulado, facilitando assim a transferência de um sistema desenvolvido ou ajustado em simulação para robôs reais.

Kestell (2000) descreve um moderno sistema de controle ativo de ruído com gerador de sinal de referência com fase bloqueada, que emprega sensores virtuais, aplicado em um avião monomotor para redução do ruído dentro da cabine, que é um ambiente extremamente ruidoso, especialmente com respeito à redução de som em baixa frequência. Resumindo os prós e os contras da abordagem, afirma: “em geral o ‘sensor virtual de densidade de energia’ supera em desempenho o sensor real de densidade de energia; contudo, os algoritmos do sensor virtual se mostram sensíveis (em graus variados) a variações de pressão acústica de pequeno comprimento de onda, dos campos de som primários e secundários”.

Oosterom e Babuška (2000) descrevem um sensor virtual para aceleração de aeronave usada em um sistema de detecção e identificação de falha, usando métodos de lógica nebulosa, em que o sensor virtual pode ser aplicado, dentro do sistema de

monitoramento, em lugar de um dos vários sensores redundantes. Esta abordagem auxilia na identificação de sensores reais defeituosos.

Haley relata um “sensor virtual” para medir a quantidade de *Listeria monocytogenes* que contamina produtos processados de carne. Baseado em um modelo fenomenológico que descreve o índice cinético de mortalidade térmica para o organismo em um tipo específico de carne processada e um modelo numérico de transferência de calor para pasteurização, o modelo retorna uma predição da quantidade de organismos residuais (SIEGEL, 2001).

Um dos parâmetros mais importantes dos processos em biotecnologia são as concentrações de biomassa. A biomassa é comumente definida como a massa ou o número de células, ou como o total da massa microbial, apesar de que a quantidade exata de interesse varia com a aplicação. A determinação exata, contínua e *on-line* da concentração de biomassa é um dos principais desejos em biotecnologia. O método mais comum para determinar a concentração de biomassa é extrair amostras para gravimetria (medição do peso seco da célula), ou também o ensaio de espectrofotometria (densidade óptica) durante a fermentação. Nos processos biotecnológicos, é pesquisada a produção de altas concentrações de biomassa e/ou produto secundário. Estes processos biológicos são afetados por muitas variáveis que têm uma influência direta no metabolismo das células, tais como temperatura, pH, taxa de evolução do oxigênio, taxa de evolução do dióxido de carbono e o tempo transcorrido na fermentação. Para otimizar a produção, é necessário controlar as variáveis mais importantes do processo. Em um monitoramento tradicional de biomassa e produto, se o desempenho da fermentação for ruim, a ação para correção da fermentação poderá resultar tardia e desastrosa. Ascencio e Herrera (1998); Ascencio e Galicia (2000) e González e Ascencio (2000) propõem o uso de redes neurais artificiais para a estimação inferencial, ou seja, o uso de sensores virtuais para se obter uma estimação da biomassa a partir de conjuntos diversos de medições, por exemplo, taxa de evolução de dióxido de carbono mais tempo transcorrido, e pH mais o tempo transcorrido e taxa de evolução de carbono mais taxa de evolução de oxigênio mais o tempo transcorrido.

Em produção de ferro e aços complexos, unidades de operação, incluindo alto-fornos e fornos de coqueamento, produzem gases como subprodutos. Algumas plantas antigas descarregavam estes gases na atmosfera, desperdiçando valiosa energia e causando severa poluição do ar. Uma planta de gás mistura estes gases para produzir o combustível para o forno de fundição de metal e laminadores de metal. A qualidade do gás misturado é medida por sua massa térmica. Gases com massa térmica inconsistente podem causar problemas no controle, na qualidade, e na produção, devido ao aumento ou diminuição do aquecimento. Mesmo durante produção normal, o fornecimento e a demanda de gás podem variar aleatoriamente. Unidades de operação como alto-fornos e fornos de reaquecimento podem trabalhar periodicamente causando enorme distúrbio no fluxo de gás, de pressão e na massa térmica. Analisadores de massa térmica *on-line* são disponíveis, mas não comumente usados durante a operação normal, porque são difíceis de manter e excessivamente caros. O objetivo é monitorar e controlar a massa térmica de gás automaticamente em todas as condições de operação. A empresa General Cybernation Group Inc. oferece uma solução para a medição e controle de massa térmica, usando uma tecnologia de sensor virtual, chamada CyboSoft, baseada em uma rede neural perceptron de multicamadas, para aplicação em processos de mistura de gás onde a massa térmica pode ser calculada corretamente e em tempo real (CYBOSOFT, 2003).

Um outro exemplo de aplicação surge na indústria de papel, que usa sensores virtuais para determinar parâmetros de qualidade considerados críticos para a folha de papel. Antes dos sensores virtuais, tais medidas tinham que ser feitas depois do papel ser fabricado, quando uma amostra era levada para uma análise em laboratório. Esta aplicação foi desenvolvida pela empresa Pavilion Technologies, utilizando também redes neurais artificiais, para predição das propriedades de qualidade da folha durante a fabricação do papel (MERRITT, 2002).

A empresa Pavilion Technologies produz também o que eles chamam tecnologia PEMS (sistema preditivo de monitoramento de efluente), ou seja, um sensor virtual para aplicação em plantas de tratamento de águas, com o intuito de prever a demanda bioquímica de oxigênio para controle ambiental, com aplicações em tratamento de

água em indústrias de papel, refinarias, processos químicos e estações de tratamentos de águas e resíduos urbanos (SOFT SENSOR – WASTEWATER, 2002).

Uma aplicação industrial bastante interessante é a descrita por Wil Chin que explica sobre um problema a ser considerado em instrumentação, encontrando na indústria a denominação de “parâmetros intangíveis”, da seguinte maneira: “Parâmetros intangíveis podem ser propriedades muito complexas de uma substância, tais como sabor ou gosto, cremosidade, cor, maciez ou suavidade, que podem somente ser subjetivamente definidos. Embora um parâmetro intangível possa ser vinculado à uma propriedade física da substância, não há uma definição conhecida”, e completa: “tipicamente, parâmetros intangíveis devem ser medidos sem se ter qualquer conhecimento detalhado do parâmetro” (MERRITT, 2002). Parâmetros intangíveis podem ser encontrados em todo o segmento industrial, incluindo as indústrias de óleo e gás. A indústria de analisadores virtuais ou sensores virtuais trabalha para fazer análises “impossíveis”, ou seja, sua metodologia é medir tudo que afeta a variável do processo, aplicar experiência e regras, e calcular o valor da variável sem obter uma medida física; em outras palavras, os sensores virtuais deduzem o valor.

Os parágrafos anteriores provêm uma amostragem da gama de abordagens e aplicações dos sensores virtuais, proporcionando uma visão clara de que, nos seminários da comunidade VIMS, poucos trabalhos foram e ainda estão sendo apresentados sobre este tema, mesmo sendo um dos fóruns mais adequados para fazê-lo, e que a maior parte dos trabalhos são acadêmicos, com algumas aplicações em escala industrial.

2.6 Desenvolvimento de sensores virtuais

Para implementar sensores virtuais, usaram-se algoritmos como filtros de Kalman e, mais recentemente, redes neurais, algoritmos genéticos e, também, computação nebulosa, como mostra o grande número de artigos recentes sobre o assunto.

Como descrito, existem várias abordagens diferentes para modelagem de sistemas não lineares complexos, as quais podem ser agrupadas em métodos globais e

métodos locais. Os métodos globais descrevem o sistema em estudo usando relações funcionais não lineares entre as variáveis do sistema. Como exemplos, temos os modelos em espaço de estados não lineares ou modelos caixa-preta de entrada-saída, como a popular estrutura NARX (*Nonlinear AutoRegressive with eXogenous input*) usada frequentemente em conexão com redes neurais ou wavelets (JOHANSEN; FOSS, 1993). As abordagens locais, por outro lado, tratam a complexidade e a não-linearidade de sistemas através da decomposição do problema de modelagem em um número mais simples, em muitos casos, de sub-problemas lineares. Por isso, procura-se estudar estes métodos que são conceitualmente simples e intuitivamente atraentes, pois são considerados próximos à maneira como os seres humanos resolvem problemas. Os modelos locais são normalmente mais facilmente interpretáveis que os complicados modelos globais, e as técnicas de modelagem baseadas em conjuntos nebulosos podem ser vistas como métodos de modelagem local (BANERJEE et al., 1995).

Neste trabalho, analisa-se a principal distinção que pode ser feita entre submodelos, como os do tipo Mamdani (linguística) ou modelos relacionais e os submodelos lineares locais representados por regras do tipo Takagi-Sugeno (TS). Levando em consideração que o objetivo é o uso de técnicas de identificação para a construção de modelos nebulosos de dados medidos, a obtenção de um modelo nebuloso baseado em regras usando esses dados requer uma intensa tarefa de identificação, como foi descrita no item anterior.

Algumas destas tarefas podem ser formuladas como um problema de otimização não-linear, para o qual um número de técnicas têm sido propostas, tais como métodos de aprendizado neural (JANG, 1992; OLIVEIRA, 1993), mínimos quadrados ortogonal (WANG, 1994), aprendizado indutivo (ROSS, 2004), raciocínio comprobatório (BALDWIN et al., 1995), ou agrupamento (*clustering*) nebuloso (BABUŠKA; VERBRUGGEN, 1995; KAYMAK; BABUŠKA, 1995).

De forma resumida, serão analisadas algumas fases do procedimento de identificação, como a coleta dos dados, a seleção da estrutura, o agrupamento e a escolha do número de agrupamentos, a derivação do modelo nebuloso, e a simplificação e validação do modelo, que podem ser observadas no esquema

representado na Figura 11, onde se emprega a abordagem de identificação baseada em agrupamentos nebulosos (BABUŠKA; VERBRUGGEN, 1997).

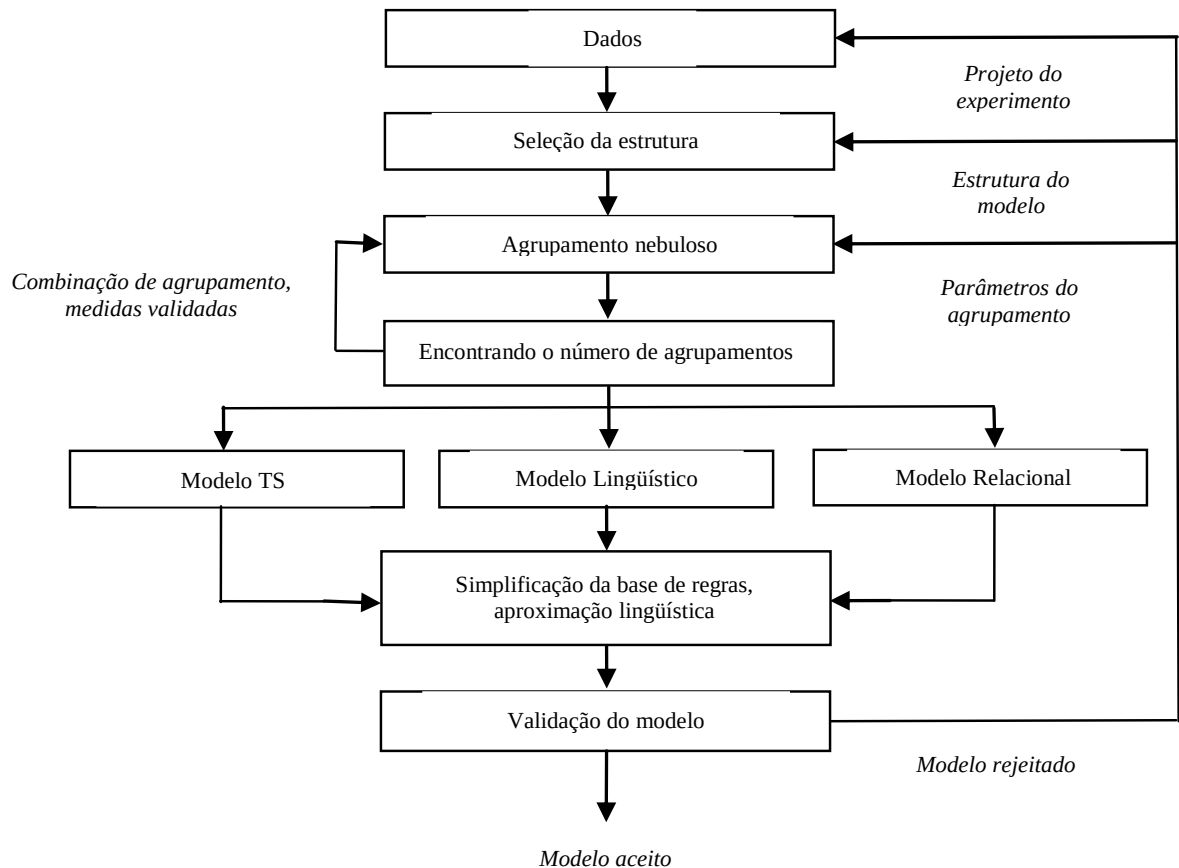


Figura 11 - Abordagem de identificação baseada em agrupamento nebuloso

Por último, mas não menos importante, e talvez o item mais significativo deste trabalho, foi o desenvolvimento de algoritmos que possibilitaram a identificação, ou seja, estimação dos parâmetros do modelo linear na parte do consequente das regras com o intuito de poder identificar parâmetros variantes no tempo, o qual tem como base o trabalho desenvolvido por Wang e Langari (1996b) em conjunto com a abordagem de identificação baseada em agrupamentos nebulosos, descrita anteriormente, proposta por Babuška e Verbruggen (1997).

3 METODOLOGIA PROPOSTA NESTE TRABALHO

O desenvolvimento de modelos matemáticos de sistemas reais é um tema importante em muitas disciplinas da engenharia e ciências. Modelos podem ser usados para simulações, análises do comportamento, melhor entendimento dos mecanismos fundamentais dos sistemas, projeto de novos processos, ou projeto de controladores.

Tradicionalmente, a modelagem é vista como uma conjunção de um completo conhecimento da natureza e do comportamento do sistema, e de um tratamento matemático apropriado que conduz a um modelo utilizável. Esta abordagem é comumente denominada modelagem do tipo caixa-branca (*white-box*). Contudo, na prática, a condição essencial para um bom entendimento dos fundamentos físicos do problema à mão demonstra ser um severo fator limitante, principalmente quando são considerados sistemas complexos. Dificuldades encontradas em modelagem convencional do tipo caixa-branca podem surgir, por exemplo, da pobre compreensão dos fenômenos básicos, valores incorretos de vários parâmetros do processo, ou da complexidade do modelo resultante. Um completo entendimento dos mecanismos básicos é praticamente impossível para a maioria dos sistemas reais. Assim, reunir um aceitável grau de conhecimento necessário para a modelagem física pode ser uma tarefa muito difícil, consumir muito tempo ou recursos ou até ser impossível. Mesmo se a estrutura do modelo estiver determinada, permanece um outro problema que é a obtenção dos valores exatos para os parâmetros, ou seja, a tarefa de identificação para estimar os parâmetros dos dados medidos no sistema. Métodos de identificação estão atualmente em um nível amadurecido somente para sistemas lineares. Muitos processos reais são, contudo, não-lineares e podem ser aproximados por modelos lineares apenas localmente.

Uma abordagem diferente assume que o processo a ser estudado possa ser aproximado através de alguma estrutura do tipo caixa-preta (*black-box*), bastante comum, usada como um aproximador de funções genéricas. O problema de modelagem reduz-se então ao requisito de uma estrutura apropriada do aproximador, a fim de se apreender corretamente as dinâmicas e não-linearidades para o sistema. Na modelagem caixa-preta, a estrutura do modelo está mal relacionada à estrutura do

sistema real, assim o problema de identificação consiste na estimação dos parâmetros do modelo. Se dados representativos do processo estiverem disponíveis, modelos do tipo caixa-preta podem ser desenvolvidos muito facilmente, sem necessitar de conhecimento específico do processo. Uma grave desvantagem desta abordagem é que a estrutura e os parâmetros destes modelos comumente não têm qualquer significado físico. Tais modelos não podem ser usados para analisar o comportamento do sistema a não ser em simulação numérica, e o modelo obtido não pode ser ampliado ou reduzido (em proporção) quando se passar a operar em outra faixa de operação do processo.

O fato de que humanos são frequentemente aptos a conduzir tarefas complexas sob considerável incerteza tem estimulado a pesquisa por padrões alternativos de modelagem e controle. Têm sido introduzidas as assim chamadas metodologias de modelagem e controle "inteligentes", que empregam técnicas motivadas por sistemas biológicos e inteligência humana para desenvolver modelos e controladores para sistemas dinâmicos. Estas técnicas exploram métodos de representação alternativa, usando, por exemplo, linguagem natural, regras, modelos qualitativos, e possuem métodos formais para incorporar informações suplementares relevantes. Modelagem e controle nebulosos são exemplos típicos de técnicas que fazem uso do conhecimento humano e processos dedutivos. Redes neurais artificiais, por outro lado, possuem capacidade de aprendizado e adaptação pela imitação do funcionamento de sistemas neurais biológicos em um nível simplificado (BABUŠKA, 2000).

3.1 Sistemas Nebulosos

Um sistema estático ou dinâmico que faz o uso de conjuntos nebulosos ou lógica nebulosa e da estrutura matemática correspondente é chamado sistema nebuloso. Existem várias maneiras de conjuntos nebulosos serem envolvidos em um sistema, tais como:

- Na descrição do sistema. Um sistema pode ser definido, por exemplo, como uma coleção de regras do tipo se-então com atributos nebulosos, ou como uma relação nebulosa. Um exemplo de uma regra nebulosa descrevendo a relação

entre a potência de aquecimento e a variação da temperatura em uma sala pode ser:

*Se a potência de aquecimento é alta **então** a temperatura aumentará rapidamente.*

- Na especificação dos parâmetros do sistema. O sistema pode ser definido por uma equação algébrica ou diferencial, na qual os parâmetros são números nebulosos em lugar de números reais. Como exemplo considere uma equação: $y = \tilde{3}x_1 + \tilde{5}x_2$, onde $\tilde{3}$ e $\tilde{5}$ são números nebulosos “cerca de 3” e “cerca de 5”, respectivamente, definidos pelas funções de pertinência. Números nebulosos expressam a incerteza nos valores dos parâmetros.
- A entrada, saída e as variáveis de estado de um sistema podem ser conjuntos nebulosos. Entradas nebulosas podem ser lidas de sensores com muito ruído (dados ruidosos), ou então serem quantidades relacionadas à percepção humana, tais como conforto, beleza, dentre outros. Sistemas nebulosos podem processar tal informação, que não é o caso com sistemas convencionais.

Um sistema nebuloso pode simultaneamente ter vários dos atributos anteriores. A Tabela 1 apresenta uma visão geral das relações entre descrições de sistemas e variáveis nebulosas e *crisp* (BABUŠKA, 2000).

Tabela 1 - Informação *crisp* e nebulosa em sistemas

descrição do sistema	dados de entrada	dados de saída resultante	estrutura matemática
<i>crisp</i>	<i>crisp</i>	<i>crisp</i>	análise funcional, álgebra linear, dentre outros.
<i>crisp</i>	nebuloso	nebuloso	princípio da extensão (uma função <i>crisp</i> pode ser avaliada por um argumento nebuloso (ZADEH apud BABUŠKA, 2000))
nebuloso	<i>crisp</i> /nebuloso	nebuloso	cálculo relacional nebuloso, inferência nebulosa

Os mais comuns são os sistemas nebulosos definidos por meio de regras se-então: sistemas nebulosos baseados em regras, que para muitos autores podem por simplicidade serem chamados modelos nebulosos. Portanto, sistemas nebulosos podem servir para diferentes fins, tais como modelagem, análise de dados, predição ou controle.

3.2 Relevância da modelagem nebulosa

Conhecimento incompleto ou vago a respeito do sistema. A teoria convencional de sistemas conta com modelos matemáticos exatos (*crisp*) de sistemas, tais como equações algébricas e equações diferenciais ou de diferenças. Para alguns sistemas, tais como sistemas eletro-mecânicos, modelos matemáticos podem ser obtidos. Isto ocorre devido às leis físicas que governam os sistemas serem bem entendidas. Para um grande número de problemas reais, contudo, a reunião de um satisfatório grau de conhecimento necessário para a modelagem física é uma tarefa difícil ou impossível. Na maioria dos sistemas, os fenômenos básicos são conhecidos apenas parcialmente e modelos matemáticos exatos não podem ser derivados ou são muito complexos para serem úteis. Exemplos de tais sistemas podem ser encontrados nas indústrias químicas ou de alimentos, biotecnologia, ecologia, finanças, sociologia, dentre outras. Uma porção importante de informação a respeito destes sistemas é disponível através do conhecimento de especialistas humanos, operadores e projetistas de processo. Este conhecimento pode ser vago e incerto para ser expresso por funções matemáticas. É, contudo, frequentemente possível descrever o funcionamento dos sistemas por meio de uma linguagem natural, na forma de regras do tipo se-então, e sistemas baseados em regras nebulosas podem ser usados como modelos baseados em conhecimento construídos usando conhecimento de especialistas em um dado campo de interesse (PEDRYCZ, 1990). Por isso, sistemas nebulosos são similares a sistemas especialistas e estudados amplamente em inteligência artificial “simbólica” (BABUŠKA, 2000).

Processamento adequado de informação imprecisa. Computação numérica precisa com modelos matemáticos convencionais somente faz sentido quando os parâmetros e os

dados de entrada são rigorosamente conhecidos. Como este não é frequentemente o caso, uma estrutura de modelagem é necessária, a qual possa adequadamente processar não somente as informações dadas, mas também as incertezas associadas. A abordagem estocástica é uma forma tradicional de tratar com incertezas, porém, tem sido observado que nem todo tipo de incerteza pode ser tratada através dela. Várias abordagens alternativas têm sido propostas, sendo uma delas a lógica nebulosa.

Modelagem e identificação transparente (caixa-cinza). Identificação de sistemas dinâmicos de medidas de entrada e saída é um importante tema de pesquisa científica com vasto alcance de aplicações práticas. Muitos sistemas reais são inerentemente não-lineares e não podem ser representados por modelos lineares usados em identificação convencional de sistemas (LJUNG, 1999). Recentemente, há um grande destaque no desenvolvimento de métodos para identificação de sistemas não-lineares de dados medidos. Redes neurais artificiais e modelos nebulosos fazem parte das mais populares estruturas de modelo utilizadas. Do ponto de vista de entrada e saída, sistemas nebulosos são funções matemáticas flexíveis que podem aproximar outras funções ou dados (medidas) com uma exatidão desejada. Comparada a outras técnicas de aproximação conhecidas, tais como redes neurais artificiais, os sistemas nebulosos proporcionam uma representação mais transparente do sistema em estudo, que é devido principalmente à possível interpretação linguística na forma de regras. A estrutura lógica das regras facilita o entendimento e análise do modelo em uma forma semiquantitativa, próximo à forma do raciocínio humano sobre o mundo real (BABUŠKA, 2000).

Modelos de sistemas nebulosos formam uma classe especial de modelos de sistemas que usam os mecanismos da lógica nebulosa para representar as características essenciais de um sistema. De um ponto de vista formal, sistemas nebulosos podem ser considerados como uma alternativa aos paradigmas de modelagem linear, não-linear e neural. Entretanto, modelos de sistemas nebulosos possuem uma característica singular que não está disponível em muitos outros tipos de técnicas de modelagem formal, ou seja, a habilidade de imitar os mecanismos de

raciocínio aproximado realizados na mente humana. Os mais comuns modelos de sistemas nebulosos consistem em coleções de regras lógicas com atributos incertos, estas regras junto com o mecanismo de raciocínio são o núcleo do modelo nebuloso (NGUYEN; SUGENO, 1998).

3.3 Modelos nebulosos baseados em regras

Em sistemas nebulosos baseados em regras, as relações entre as variáveis são representadas por meio de regras nebulosas do tipo se-então da seguinte forma geral:

$$\text{Se proposição antecedente **então** proposição consequente.} \quad (8)$$

A proposição antecedente é sempre uma proposição nebulosa do tipo “ $\tilde{x}$ é A ” onde $\tilde{x}$ é uma variável linguística e A é uma constante linguística (termo). O valor verdadeiro da proposição (um número real entre zero e um) depende do grau de equiparação (similaridade) entre $\tilde{x}$ e A . Dependendo da forma do consequente, dois tipos principais de modelos nebulosos baseados em regras são diferenciados:

- *O modelo nebuloso linguístico ou tipo Mamdani:* ambos o antecedente e o consequente são proposições nebulosas.
- *O modelo nebuloso tipo Takagi-Sugeno (TS):* o antecedente é uma proposição nebulosa e o consequente é uma função exata.

3.3.1 Modelo nebuloso linguístico

O modelo nebuloso linguístico foi introduzido como um método para capturar o conhecimento (semi-)qualitativo disponível na forma de regras se-então:

$$\mathcal{R}_i: \text{Se } \tilde{x} \text{ é } A_i \text{ então } \tilde{y} \text{ é } B_i, \quad i = 1, 2, \dots, K \quad (9)$$

Aqui $\tilde{x}$ é a variável linguística de entrada (antecedente) e A_i são os termos linguísticos antecedentes (constantes). Similarmente, $\tilde{y}$ é a variável linguística de saída (consequente) e B_i são os termos linguísticos consequentes. Os valores de $\tilde{x}$ ($\tilde{y}$) e os termos linguísticos A_i (B_i) são conjuntos nebulosos definidos nos domínios: $x \in X \subset \mathbb{R}^p$ e $y \in Y \subset \mathbb{R}^q$.

As funções de pertinência dos conjuntos nebulosos antecedentes (consequentes) são então os mapeamentos: $\mu(x): X \rightarrow [0,1]$, $\mu(y): Y \rightarrow [0,1]$.

Conjuntos nebulosos A_i definem regiões nebulosas no espaço antecedente, para os quais as respectivas proposições consequentes são válidas. Os termos linguísticos A_i e B_i são comumente selecionados dos conjuntos de termos predefinidos, tais como *Pequeno*, *Médio*, dentre outros.

Designando-se estes conjuntos por $\mathcal{A}$ e $\mathcal{B}$, respectivamente, tem-se $A_i \in \mathcal{A}$ e $B_i \in \mathcal{B}$. A base de regras $\mathcal{R} = \{\mathcal{R}_i | i = 1, 2, \dots, K\}$ e os conjuntos $\mathcal{A}$ e $\mathcal{B}$ constituem a base de conhecimento do modelo linguístico.

A fim de se poder usar o modelo linguístico, é necessário um algoritmo que permita calcular o valor de saída, dado algum valor de entrada. Este algoritmo é chamado mecanismo de inferência nebulosa e para o modelo linguístico, o mecanismo de inferência pode ser derivado, por exemplo, usando cálculo relacional nebuloso (BABUŠKA, 2000).

3.3.2 Modelo nebuloso tipo Takagi-Sugeno

O modelo linguístico, como visto no item anterior, descreve um dado sistema por meio de regras linguísticas do tipo se-então com proposição nebulosa no antecedente assim como no consequente. O modelo nebuloso do tipo Takagi-Sugeno (TS) (TAKAGI; SUGENO, 1985), por outro lado, usa funções exatas nos consequentes. Por esta razão, ele pode ser visto como uma combinação de modelagem linguística e regressão matemática, no sentido que os antecedentes descrevem regiões nebulosas no

espaço de entrada, no qual funções consequentes são válidas. As regras do tipo TS têm a seguinte forma:

$$\mathcal{R}_i: \text{Se } \mathbf{x} \text{ é } A_i \text{ então } y_i = f_i(\mathbf{x}), \quad i = 1, 2, \dots, K \quad (10)$$

Contrariamente ao modelo linguístico, a entrada $\mathbf{x}$ é uma variável exata (entradas linguísticas são em princípio possíveis, mas exigiria o uso do princípio da extensão de Zadeh para calcular o valor nebuloso de y_i). As funções f_i são tipicamente da mesma estrutura, somente os parâmetros em cada regra são diferentes. Geralmente, f_i é um vetor com valor de função, mas para facilitar a notação considera-se um escalar f_i na sequência. Uma parametrização simples e útil é a forma “*affine*” (linear em parâmetros), produzindo as regras:

$$\mathcal{R}_i: \text{Se } \mathbf{x} \text{ é } A_i \text{ então } y_i = \mathbf{a}_i^T \mathbf{x} + b_i, \quad i = 1, 2, \dots, K \quad (11)$$

onde $\mathbf{a}_i$ é um vetor de parâmetros e b_i é um *offset* escalar. Este modelo é chamado um *modelo do tipo TS affine* (BABUŠKA, 2000).

Neste trabalho serão usados modelos nebulosos do tipo Takagi-Sugeno, pois combinam uma descrição global baseada em regras e modelos lineares locais. Embora menos genéricos que os modelos do tipo Mandani, são mais fáceis de serem gerados e implementados, sendo que a cada região do espaço está associada uma equação linear, cujos parâmetros podem ser calculados pelo método dos mínimos quadrados. As regras deste modelo podem ser de duas formas, de acordo com a concepção do modelo.

Na primeira forma, cada variável de entrada é tratada separadamente, cada qual podendo ser associada aos seus respectivos conjuntos. Por exemplo, imaginando-se um sistema que possua duas variáveis de entrada u_1 e u_2 , onde um valor de u_1 pode ser classificado como pertencente, com diferentes graus, aos conjuntos A_{1i} e um valor de u_2 que pode ser associado aos conjuntos A_{2i} , as regras nebulosas assumiriam o formato:

$$R_x : \text{Se } (u_1(k) \in A_{1i} \text{ E } u_2(k) \in A_{2j}) \text{ então } y_x(k) = b_{x0} + b_{x1}u_1(k) + b_{x2}u_2(k) \quad (12)$$

Na segunda forma, os valores das variáveis são tratados agrupados na forma de um ponto ou vetor. Os conjuntos A_i , agora, são regiões definidas no espaço de dimensão p , onde p é o número de variáveis de entrada. Neste modelo, as regras assumem o formato (BABUŠKA, 2000) e (OLIVEIRA; GARCIA, 2003):

$$R_x : \text{Se } (z(k) \in A_x) \text{ então } y_x(k) = b_{x0} + b_{x1}u_1(k) + b_{x2}u_2(k) + \dots + b_{xp}u_p(k) \quad (13)$$

onde $z(k) = (u_1(k), u_2(k), \dots, u_p(k))$.

3.4 Construção de modelos nebulosos

Tradicionalmente, a lógica nebulosa tem sido vista pela comunidade de inteligência artificial como uma abordagem para tratar incertezas. Nos anos 90, contudo, a lógica nebulosa surgiu como um paradigma para aproximação de um mapeamento funcional (BABUŠKA, 2000). Esta visão moderna e complementar sobre a tecnologia oferece novas percepções sobre os fundamentos da lógica nebulosa, assim como novos desafios relativos à identificação de modelos nebulosos.

A identificação de modelos nebulosos consiste de três sub-problemas básicos: a identificação de estrutura, a estimação de parâmetros e a validação do modelo (BABUŠKA, 2000).

Identificação de estrutura: envolve a descoberta das variáveis de entrada relevantes dentre todas as possíveis variáveis de entrada, envolvendo ainda a especificação das funções de pertinência, o particionamento do espaço de entrada e a determinação do número de regras nebulosas compreendendo o modelo básico.

Estimação de parâmetros: envolve a determinação dos parâmetros desconhecidos no modelo usando algum método de otimização baseado tanto em informação linguística obtida de especialistas humanos como em dados numéricos obtidos do sistema físico real. A especificação da estrutura e a estimação de parâmetros estão ligadas, e cada uma delas não pode ser independentemente identificada sem fazer uso do outro.

Validação do modelo: envolve o teste do modelo baseado em algum critério de desempenho (por exemplo, precisão).

Se o modelo não passar no teste de validação, sua estrutura deve ser modificada e re-estimados os seus parâmetros, repetindo-se este processo até que um modelo satisfatório seja encontrado.

Devem ser observadas ainda duas das mais importantes questões em modelagem nebulosa de problemas de dimensão elevada, ou seja, de problemas complexos:

- 1) identificação da partição nebulosa do espaço de entrada, e
- 2) tratamento da explosão do número de regras, ou seja, o aumento exponencial do número de regras quando o número de variáveis de entrada aumenta.

Partições nebulosas do espaço de entrada definem o antecedente de um modelo nebuloso. O número de partições nebulosas usualmente determina o número de regras nebulosas incluindo também o modelo básico. Existem na literatura três tipos de partições nebulosas: 1) a partição em grade, 2) a partição em árvore e 3) a partição em dispersão (YEN, 1999).

A obtenção do modelo nebuloso ou identificação nebulosa pode ser feita de três modos. O primeiro é subjetivo, pois utiliza o conhecimento prévio de um especialista. O segundo é baseado em modelos estatísticos. E o terceiro modo usa a “clusterização” como forma de determinação de agrupamentos de dados. Neste trabalho é empregada a “clusterização”, pois se deseja utilizar uma forma automática de geração de modelos.

Assim, a partir de um conjunto de dados recolhidos da planta (dados de treinamento), é feita a “clusterização” por um dos métodos descritos adiante. Com os conjuntos nebulosos já definidos, são geradas as regras, e, a partir de então, os parâmetros dos consequentes são estimados. Obtém-se, dessa maneira, um modelo da planta estudada.

3.5 Clusterização ou agrupamento nebuloso

Como já visto nos itens anteriores, modelos nebulosos descrevem os sistemas estabelecendo relações entre as variáveis relevantes do sistema (por exemplo, entradas e saídas) na forma de regras do tipo se-então. Cada regra mapeia uma região nebulosa do espaço de premissa para outra região no espaço consequente. Inferência nebulosa

assegura a interpolação entre as regras. Conceber um bem sucedido modelo nebuloso de um sistema dinâmico requer o ajuste de um número de parâmetros, como a escolha das variáveis relevantes, o número de funções de pertinência por variável, a forma e os parâmetros da função de pertinência, o número de regras na base de regras, dentre outros. A sintonização heurística destas variáveis é muito tediosa e consome muito tempo, e a geração automática de uma base de regras das medidas do sistema é portanto altamente desejável.

Um dos métodos que pode ser aplicado para a geração automática de um modelo nebuloso é o agrupamento (*clustering*) nebuloso. Algoritmos de agrupamentos nebulosos são não supervisionados e formam uma partição nebulosa de pontos de dados em um dado número de grupos (*clusters*). Aplicando-se agrupamento nebuloso nos dados obtidos das medições feitas no sistema, um modelo nebuloso do sistema pode ser obtido. Antes da aplicação de um algoritmo de agrupamento nebuloso, a estrutura do problema deve ser determinada, identificando-se as variáveis relevantes do sistema. Então, o número de agrupamentos requeridos deve ser selecionado. Este número determinará o número de regras na base de regras resultante (KAYMAK; BABUŠKA, 1995).

Agrupamentos de dados numéricos formam a base de muitos algoritmos de classificação e modelagem de sistemas. O propósito do *clustering* é identificar agrupamentos naturais de dados de um grande conjunto de dados para produzir uma representação concisa de um comportamento do sistema.

Métodos de identificação baseados em agrupamentos nebulosos se originam de análises de dados e reconhecimento de padrões, onde o conceito de classificação (graduação) de pertinência é usado para representar o grau para o qual um dado objeto, representado como um vetor de características, é similar a algum objeto padrão. O grau de similaridade pode ser calculado usando-se uma medida de distância satisfatória. Baseado na similaridade, vetores característicos podem ser agrupados tal que os vetores dentro de um agrupamento sejam tão similares (próximos) quanto possível, e vetores de diferentes agrupamentos sejam tão diferentes quanto possível.

Esta idéia de agrupamento nebuloso é ilustrada pela Figura 12, onde o dado é agrupado em dois grupos com padrões v_1 e v_2 , usando-se a medida de distância

Euclidiana. A partição dos dados é expressa na *matriz de partição nebulosa* cujos elementos μ_{ij} são graus de pertinência dos pontos de dados $[x_i, y_i]$ em um agrupamento nebuloso com padrões v_j .

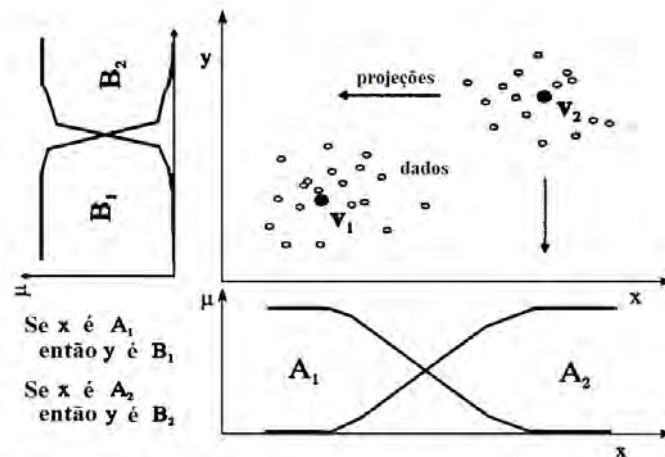


Figura 12 - Interpretação baseada em regra de agrupamentos nebulosos

Regras do tipo se-então nebulosas podem ser extraídas através da projeção dos agrupamentos sobre os eixos. A Figura 12 mostra um conjunto de dados com dois agrupamentos aparentes e duas regras nebulosas associadas. O conceito de similaridade dos dados para um padrão dado deixa suficiente espaço para a escolha de uma medida apropriada de distância e do próprio caráter do padrão.

Por exemplo, os padrões podem ser definidos como subespaços lineares (BABUŠKA, 2000), ou os agrupamentos podem ser elipsóides com formas determinadas adaptativamente (GUSTAFSON; KESSEL, 1978), como na Figura 13.

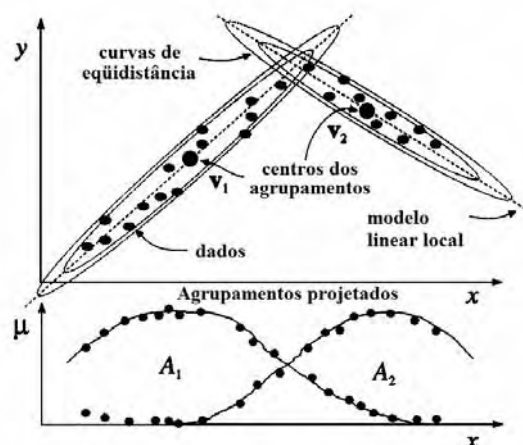


Figura 13 - Agrupamentos nebulosos hiper-elipsoidais

De tais agrupamentos, as funções de pertinência dos antecedentes e os parâmetros dos consequentes do modelo tipo Takagi-Sugeno podem ser extraídos:

$$\begin{aligned} \text{Se } x \text{ é } A_1 \text{ então } y &= a_1x + b_1, \\ \text{Se } x \text{ é } A_2 \text{ então } y &= a_2x + b_2. \end{aligned} \quad (14)$$

Cada agrupamento obtido é representado por uma regra no modelo tipo Takagi-Sugeno. As funções de pertinência para os conjuntos nebulosos A_1 e A_2 são geradas por projeções com relação ao ponto da matriz de partição sobre as variáveis antecedentes. Estes conjuntos nebulosos definidos com relação ao ponto são então aproximados por uma função paramétrica apropriada. Os parâmetros consequentes para cada regra são estimados usando-se mínimos quadrados (BABUŠKA, 2000).

Assim, a clusterização é o primeiro passo para que seja gerado o núcleo de inferência nebuloso a partir de dados coletados do processo e, com o resultado da clusterização, é possível se determinarem as funções de pertinência, os graus de ativação e os parâmetros dos consequentes das regras.

3.5.1 *Fuzzy C-Means* (FCM)

A técnica de agrupamento de dados em que cada ponto pertence a um grupo com um certo grau, que é especificado por uma graduação de pertinência, é chamada de *Fuzzy C-means*. Esta técnica foi originalmente introduzida por James Bezdek (apud BABUŠKA, 2000) como um aperfeiçoamento nos métodos de agrupamentos iniciais. Tal técnica proporciona um método para agrupar pontos de dados que povoam algum espaço multidimensional em um número específico de agrupamentos diferentes.

Em alguns dos métodos utilizados para obtenção de modelos nebulosos, o FCM foi empregado na clusterização dos dados de entrada (dados de treinamento). O FCM procura por regiões do espaço de entrada onde haja agrupamentos de pontos, determinando seus centros. Neste algoritmo de clusterização, a pertinência de um ponto a um agrupamento decresce com a distância e independe da direção, podendo-se dizer que, neste caso, os agrupamentos são esféricos.

O algoritmo FCM tem como objetivo minimizar a função:

$$J = \sum_{i=1}^c \sum_{k=1}^n (\mu_{ik})^m \cdot \|z_k - v_i\|^2 \quad (15)$$

onde:

$z_k = [u_1(k), u_2(k), \dots, u_p(k)]$: ponto de treinamento k ;

$v_i = [x_{i1}, x_{i2}, \dots, x_{ip}]$: centro do agrupamento i ;

μ_{ik} : pertinência do ponto k ao agrupamento i , $\mu_{ik} \in [0;1]$;

c : número desejado de agrupamentos;

n : número de pontos coletados;

m : determina a “nebulosidade” dos agrupamentos, normalmente assume-se $m=2$;

p : número de variáveis de entrada.

A pertinência de um ponto a um agrupamento pode ser calculada pela equação:

$$\mu_{ik} = \frac{1}{\sum_{j=1}^c (d_{ki} / d_{kj})^{2/(m-1)}} \quad (16)$$

onde:

$d_{ki} = \|z_k - v_i\|$, utiliza-se a norma euclidiana;

$1 \leq i \leq c, \quad 1 \leq k \leq p$

Caso $d_{ki}=0$ para algum $i=s$, então $\mu_{sk}=1$; e, $\forall i \neq s, \mu_{ik} = 0$.

O conjunto de variáveis de pertinência μ_{ik} deve obedecer a duas restrições:

$$\sum_{i=1}^c \mu_{ik} = 1, \quad k = 1, \dots, n \quad (17a)$$

$$0 < \sum_{k=1}^n \mu_{ik} < n, \quad i = 1, \dots, c \quad (17b)$$

A primeira restrição visa evitar a solução trivial da função (16), isto é, evita a solução onde $\mu_{ik} = 0$ para $i = 1, \dots, c$ e $k = 1, \dots, p$. A segunda restrição impede que um agrupamento esteja totalmente vazio ou que contenha todos os pontos com grau de pertinência 1.

As coordenadas dos centros $\mathbf{v}_i$ podem ser calculadas pela equação:

$$\mathbf{v}_i = \frac{\sum_{k=1}^n (\mu_{ik})^m \cdot \mathbf{z}_k}{\sum_{k=1}^n (\mu_{ik})^m} \quad (18)$$

Inicia-se ou com uma matriz de centros $\mathbf{V}=[\mathbf{v}_i]=[x_{ij}]$, $i= 1,2,...,c$; $j= 1,2,...,p$; ou com uma matriz de partição nebulosa $\mathbf{U}=[\mu_{ik}]$, $i = 1,2,...,c$; $k = 1,2,...,n$; geradas randomicamente (obedecendo as restrições (17a) e (17b)). E, através das equações (14) e (15), podem-se calcular iterativamente os centros dos agrupamentos que minimizem a equação (13).

Dada uma tolerância ξ , o critério de parada pode ser definido como sendo:

- a. $\max(|\mu_{ik}^l - \mu_{ik}^{l-1}|) < \xi$ ou $\|\mathbf{U}^l - \mathbf{U}^{l-1}\| < \xi$, caso se tenha iniciado por uma matriz $\mathbf{U}$; ou,
- b. $\max(|x_{ij}^l - x_{ij}^{l-1}|) < \xi$ ou $\|\mathbf{V}^l - \mathbf{V}^{l-1}\| < \xi$, caso se tenha iniciado por uma matriz $\mathbf{V}$,

onde o índice l indica o valor calculado na l -ésima iteração (OLIVEIRA; GARCIA, 2003) e (WANG; LANGARI, 1996a).

3.5.2 Clusterização pelo algoritmo de Gustafson-Kessel

Como ocorre com o algoritmo FCM, este algoritmo também procura por agrupamentos de pontos, determinando seus centros. No entanto, é mais genérico, pois os agrupamentos determinados são elipsóides. Desta forma, a pertinência de um ponto a um agrupamento decresce de maneira diferente para diferentes direções. Cada agrupamento, agora, possui um diferente sistema de coordenadas onde cada eixo possui um peso diferente. O algoritmo foi extraído de (BABUŠKA; VERBRUGGEN, 1997).

O objetivo continua sendo minimizar a função:

$$J = \sum_{i=1}^c \sum_{k=1}^n (\mu_{ik})^m \cdot d^2(\mathbf{z}_k, \mathbf{v}_i) \quad (17)$$

sujeita às restrições (14a) e (14b). Mas a nova distância é calculada pela equação:

$$d^2(\mathbf{z}_k, \mathbf{v}_i) = d_{ki}^2 = (\mathbf{z}_k - \mathbf{v}_i) \cdot \mathbf{M}_i \cdot (\mathbf{z}_k - \mathbf{v}_i)^T \quad (18)$$

onde $\mathbf{M}_i$ é uma matriz positiva definida gerada de acordo com a forma de um determinado agrupamento i . A matriz $\mathbf{M}_i$ pode ser calculada pela matriz de covariância do agrupamento i , $\mathbf{F}_i$, como a seguir:

$$\mathbf{M}_i = \det(\mathbf{F}_i)^{\frac{1}{p}} \cdot \mathbf{F}_i^{-1} \quad (19)$$

onde as matrizes de covariância $\mathbf{F}_i$ são calculadas por:

$$\mathbf{F}_i = \frac{\sum_{k=1}^n (\mu_{ik})^m \cdot (\mathbf{z}_k - \mathbf{v}_i)^T \cdot (\mathbf{z}_k - \mathbf{v}_i)}{\sum_{k=1}^n (\mu_{ik})^m} \quad i = 1, 2, \dots, c \quad (20)$$

Cada iteração (l) do algoritmo de Gustafson-Kessel consiste em quatro passos. Definidos o número de agrupamentos c , a constante m e uma tolerância ξ , e, sendo dado um conjunto de dados de treinamento $Z = \{\mathbf{z}_k\}$, $k=1, 2, \dots, n$, pode-se iniciar a clusterização a partir de uma matriz de partição nebulosa inicial $\mathbf{U} = [\mu_{ik}]$, $i=1, 2, \dots, c$; $k=1, 2, \dots, n$; ou através de uma matriz de centros $\mathbf{V} = [\mathbf{v}_i] = [x_{ij}]$, $i=1, 2, \dots, c$; $j=1, 2, \dots, p$. Os passos de cada iteração são:

a) Calcular os centros dos agrupamentos:

$$\mathbf{v}_i^l = \frac{\sum_{k=1}^n (\mu_{ik}^l)^m \cdot \mathbf{z}_k}{\sum_{k=1}^n (\mu_{ik}^l)^m} \quad i = 1, 2, \dots, c \quad (21)$$

- b) Calcular as matrizes de covariância pela equação (20).
- c) Calcular as distâncias dos pontos aos centros dos agrupamentos pelas equações (18) e (19), $i=1, \dots, c$ e $k=1, \dots, n$.
- d) Atualizar a matriz de partição nebulosa $\mathbf{U}=[\mu_{ik}]$ pela equação (16), utilizando as distâncias calculadas no item c. Caso $d_{ki}=0$ para algum $i=s$, então $\mu_{sk}=1$; e, $\forall i \neq s, \mu_{ik}=0$.

Quando se inicia por uma matriz $\mathbf{U}$, a sequência de passos é **a**, **b**, **c** e **d**. Por outro lado, iniciando-se por uma matriz de centros $\mathbf{V}$, a sequência é **b**, **c**, **d** e **a**.

Dada uma tolerância ξ , o critério de parada pode ser definido como sendo:

- $\max(|\mu_{ik}^l - \mu_{ik}^{l-1}|) < \xi$ ou $\|\mathbf{U}^l - \mathbf{U}^{l-1}\| < \xi$, caso se tenha iniciado por uma matriz $\mathbf{U}$.
- $\max(|x_{ij}^l - x_{ij}^{l-1}|) < \xi$ ou $\|\mathbf{V}^l - \mathbf{V}^{l-1}\| < \xi$, caso se tenha iniciado por uma matriz $\mathbf{V}$.

onde o índice l indica o valor calculado na l -ésima iteração.

A Figura 14 ilustra qualitativamente a diferença entre os agrupamentos gerados pelo FCM e pelo algoritmo de Gustafson-Kessel (GK). As curvas vermelhas indicam pontos do espaço de iguais graus de pertinência ao respectivo agrupamento. Pode ser observada a liberdade, no algoritmo de GK, quanto à direção e pesos dos eixos dos agrupamentos. A clusterização pelo algoritmo de GK encontra muitas aplicações no processamento de imagens e reconhecimento de padrões (OLIVEIRA; GARCIA, 2003) (BABUŠKA; VERBRUGGEN, 1997).

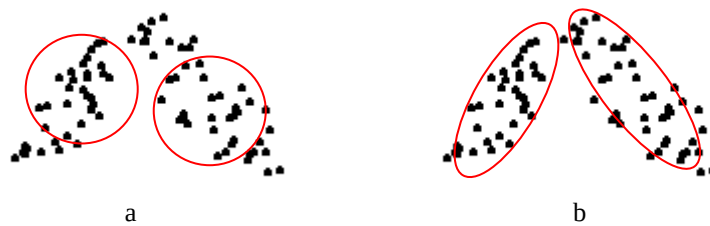


Figura 14 - a- Ilustração de agrupamentos gerados pelo FCM.
b- Ilustração de agrupamentos gerados por Gustafson-Kessel.

3.6 Estimação dos parâmetros consequentes das regras

Nos modelos Takagi-Sugeno abordados, é necessária a estimação dos parâmetros consequentes das regras. Este problema pode ser transformado num problema de estimação dos parâmetros do modelo de regressão linear cuja solução é descrita a seguir.

Dado um conjunto de pontos $(\varphi_{k1}, \varphi_{k2}, \dots, \varphi_{kp}) \rightarrow y_k, k = 1, 2, \dots, n$, pode-se organizá-los num vetor de saída $\mathbf{Y}=[y_1, y_2, \dots, y_n]^T$ e numa matriz de entradas:

$$\mathbf{\Phi} = \begin{bmatrix} \varphi_{11} & \cdots & \varphi_{1p} \\ \vdots & \ddots & \vdots \\ \varphi_{n1} & \cdots & \varphi_{np} \end{bmatrix}$$

Deseja-se determinar o vetor de parâmetros $\mathbf{B}=[b_1, b_2, \dots, b_p]^T$ tal que $\mathbf{\Phi B}=\mathbf{Y}$. Caso $n=p$, $\mathbf{\Phi}$ será uma matriz quadrada, e então, caso haja solução, $\mathbf{B}=\mathbf{\Phi}^{-1}\mathbf{Y}$.

No entanto, ao se lidar com identificação de sistemas $n \neq p$, ou mais especificamente $n > p$, deseja-se encontrar o vetor de parâmetros $\mathbf{B}$, que determine a solução que melhor aproxime os pontos dados. Pelo método dos mínimos quadrados, o objetivo será minimizar a função:

$$J(\mathbf{B}) = \|\mathbf{Y} - \mathbf{\Phi} \cdot \mathbf{B}\|^2 \quad (22)$$

Assim, o vetor $\mathbf{B}$ poderia ser calculado por $\mathbf{B}=(\mathbf{\Phi}^T \mathbf{\Phi})^{-1} \mathbf{\Phi}^T \mathbf{Y}$. A resolução direta desta equação não é simples e nem rápida. Mesmo utilizando bons programas computacionais, a resolução da inversão de matriz pode gerar diversos erros, e é comum que esta operação não possa ser efetuada devido à matriz ser singular para a precisão utilizada pelo *software*. Conforme proposto por Wang e Langari (WANG; LANGARI, 1996a), é conveniente o uso de um algoritmo recursivo composto das equações:

$$\begin{aligned}
\mathbf{f}_k^l &= (\mathbf{C}_{k-1}^l)^T \cdot \boldsymbol{\Phi}_k^T \\
\gamma_k^l &= \left[1 + (\mathbf{f}_k^l)^T \cdot \mathbf{f}_k^l \right]^{1/2} \\
\alpha_k^l &= \gamma_k^l + 1 \\
\beta_k^l &= \frac{1}{\gamma_k^l \cdot \alpha_k^l} \\
\boldsymbol{\sigma}_k^l &= \beta_k^l \cdot \mathbf{C}_{k-1}^l \cdot \mathbf{f}_k^l \\
\mathbf{B}_k^l &= \mathbf{B}_{k-1}^l + \boldsymbol{\sigma}_k^l \cdot \left[\frac{\alpha_k^l}{\gamma_k^l} \cdot (y_k - \boldsymbol{\Phi}_k \cdot \mathbf{B}_{k-1}^l) \right] \\
\mathbf{C}_k^l &= \mathbf{C}_{k-1}^l - \boldsymbol{\sigma}_k^l \cdot (\mathbf{f}_k^l)^T
\end{aligned} \tag{23}$$

onde, $\boldsymbol{\Phi}_k$ é a k -ésima linha de $\boldsymbol{\Phi}$. Inicia-se o processo por uma matriz $\mathbf{C}$, identidade, de ordem $\mathbf{p}$. A matriz $\mathbf{B}$ inicial é tal que $b_i = 1$, $1 \leq i \leq p$. E, a cada iteração l , as equações (23) são calculadas para $k=1, 2, \dots, n$. Para o início de uma nova iteração são usados os últimos valores de $\mathbf{C}$ e $\mathbf{B}$, ou seja, $\mathbf{C}_0^l = \mathbf{C}_n^{l-1}$ e $\mathbf{B}_0^l = \mathbf{B}_n^{l-1}$.

O critério de parada, dada uma tolerância ξ , pode ser definido como:

$$\|\mathbf{B}_n^l - \mathbf{B}_n^{l-1}\| / p < \xi ; \text{ou, } \max(|b_i^l - b_i^{l-1}|) < \xi, \text{ para } \mathbf{B}_n^l \text{ e } 1 \leq i \leq p.$$

3.7 Modelos nebulosos utilizados

Os modelos utilizados são do tipo Takagi-Sugeno (TS), no entanto, cada modelo possui suas particularidades.

3.7.1 Modelo de Wang-Langari

Na abordagem de Wang e Langari (WANG; LANGARI, 1996a), as variáveis de entrada são tratadas independentemente. Os dados de cada variável são agrupados, gerando-se assim c agrupamentos por variável e consequentemente c centros de agrupamentos por variável. O conjunto de regras toma a forma:

$$R_x: \text{ Se } (u_1(k) \in A_{1i} \text{ e } u_2(k) \in A_{2j} \text{ e } \dots \text{ e } u_p(k) \in A_{pk}) \text{ então } \\ y_x(k) = b_{xo} + b_{x1}u_1(k) + b_{x2}u_2(k) + \dots + b_{xp}u_p(k)$$

onde A_{ij} refere-se ao cluster j da variável u_i , $i=1,2,\dots, p$, e $j=1,2,\dots,c$. Tem-se assim c^p regras, onde p é o número de variáveis de entrada; ou seja, devemos possuir regras para todas as combinações possíveis de conjuntos.

Inicia-se o método agrupando-se um conjunto de dados recolhidos da planta (modelo) numa matriz $\mathbf{Z}$ de dados de entrada na forma:

$$\mathbf{Z} = \begin{bmatrix} u_1^1 & \cdots & u_p^1 & y^1 \\ \vdots & \ddots & \vdots & \vdots \\ u_1^n & \cdots & u_p^n & y^n \end{bmatrix}$$

onde $u_i^k = u_i(k)$ e $y^k = y(k)$. Cada coluna desta matriz, relativa às variáveis de entrada, é extraída e clusterizada pelo FCM independentemente das demais, obtendo-se assim c agrupamentos unidimensionais por variável de entrada, definidos por c centros por variável, v_{ij} , $i=1,2,\dots,p$ e $j=1,2,\dots,c$. Estes centros podem ser agrupados numa matriz:

$$\mathbf{V} = \begin{bmatrix} v_{11} & \cdots & v_{1c} \\ \vdots & \ddots & \vdots \\ v_{p1} & \cdots & v_{pc} \end{bmatrix}$$

Tendo-se, então, os centros de cada agrupamento, calcula-se a pertinência de cada ponto de treinamento a cada agrupamento pela equação:

$$\mu_{ij}^k = \frac{1}{\sum_{q=1}^c \left(\frac{u_i^k - v_{ij}}{u_i^k - v_{iq}} \right)^{2/(m-1)}} \quad (24)$$

$i = 1, 2, \dots, p \quad j = 1, 2, \dots, c \quad k = 1, 2, \dots, n$

Agora μ_{ij}^k representa a pertinência do ponto k , com relação à variável u_i , ao agrupamento j desta mesma variável. Neste caso v_{ij} é um escalar. Cada regra é determinada por uma combinação de “clusters” de cada variável. Assim uma regra R_h é identificada por um vetor $A_h = [\lambda_1, \lambda_2, \dots, \lambda_p]$, onde $\lambda_i \in \{1, 2, \dots, c\}$, tal que:

$$R_h : \begin{array}{l} \text{Se } (u_1^k \in A_{1\lambda_1} \text{ e } u_2^k \in A_{2\lambda_2} \text{ e } \dots \text{ e } u_p^k \in A_{p\lambda_p}) \\ \text{então } y_h^k = b_{h0} + b_{h1}u_1^k + b_{h2}u_2^k + \dots + b_{hp}u_p^k \end{array}$$

O grau de ativação g de cada regra para um ponto k é calculado por:

$$g_h^k = \frac{w_h^k}{\sum_{i=1}^r w_i^k} \quad h = 1, 2, \dots, r \quad (25)$$

sendo r o número total de regras, ou seja, $r = c^p$. Um w_h^k qualquer pode ser calculado por:

$$w_h^k = \mu_{1\lambda_1}^k \cdot \mu_{2\lambda_2}^k \cdot \dots \cdot \mu_{n\lambda_p}^k \quad (26)$$

A saída deste modelo nebuloso é:

$$y_f^k = \sum_{h=1}^r g_h^k \cdot y_h^k = \sum_{h=1}^r g_h^k \cdot (b_{h0} + b_{h1} \cdot u_1^k + \dots + b_{hp} \cdot u_p^k) \quad (27)$$

Para finalizar o modelo, falta apenas calcular os parâmetros consequentes b_{ij} . O objetivo é minimizar a somatória dos erros quadráticos para o modelo alimentado com os pontos coletados. Ou seja, deseja-se minimizar:

$$J = \sum_{k=1}^n (y^k - y_f^k)^2 = \sum_{k=1}^n \left(y^k - \sum_{h=1}^r g_h^k (b_{h0} + b_{h1}u_1^k + \dots + b_{hp}u_p^k) \right)^2 \quad (28)$$

Definindo-se:

$$\begin{aligned} \mathbf{Y} &= [y^1 \quad y^2 \quad \dots \quad y^n]^T \\ \mathbf{B} &= [b_{10} \quad \dots \quad b_{1p} \quad b_{20} \quad \dots \quad b_{2p} \quad \dots \quad b_{r0} \quad \dots \quad b_{rp}]^T \\ \Phi &= \begin{bmatrix} g_1^1 & \dots & g_r^1 & g_1^1 u_1^1 & \dots & g_r^1 u_1^1 & \dots & g_1^1 u_p^1 & \dots & g_r^1 u_p^1 \\ g_1^2 & \dots & g_r^2 & g_1^2 u_1^2 & \dots & g_r^2 u_1^2 & \dots & g_1^2 u_p^2 & \dots & g_r^2 u_p^2 \\ \vdots & & & \vdots & & & \ddots & \vdots & & \\ g_1^n & \dots & g_r^n & g_1^n u_1^n & \dots & g_r^n u_1^n & \dots & g_1^n u_p^n & \dots & g_r^n u_p^n \end{bmatrix} \end{aligned} \quad (29)$$

A função objetivo J , a ser minimizada, pode então ser reescrita como:

$$J = \|\mathbf{Y} - \Phi \cdot \mathbf{B}\|^2$$

Desta forma, chega-se ao mesmo problema do item 3.6, onde a solução foi descrita.

Executando-se este procedimento obtém-se um modelo nebuloso a partir de dados coletados na planta. Para um novo valor de entrada, obtém-se a saída calculando-se, para o ponto, os graus de ativação de cada regra e calculando-se a saída pela equação (27). Diferentemente do proposto por Wang e Langari (WANG; LANGARI, 1996a), na implementação utilizada, as pertinências dos pontos de treinamento não são interpoladas para o cálculo de uma nova pertinência, mas é utilizada a própria equação (24).

A Figura 15 ilustra qualitativamente a divisão do espaço de entradas entre as regras do modelo Wang-Langari para 2 agrupamentos por variável. Os limites de ativação de cada regra são nebulosos. As linhas tracejadas apenas representam as regiões onde cada regra prevalece, ou seja, onde o grau de ativação de determinada regra é maior que as demais.

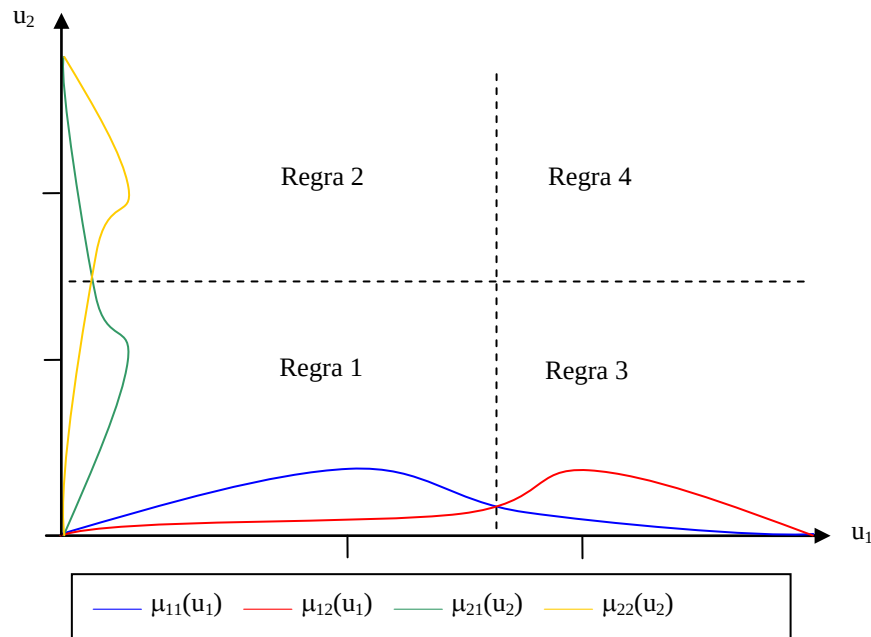


Figura 15 - Ilustração das regiões do espaço de entrada associadas às regras.

3.7.2 Modelo por “Regras Limitadas”

O uso do método descrito no item anterior leva a resultados razoáveis, no entanto, seu uso é bastante limitado. Para um número pequeno de variáveis (p) e “clusters” por variável (c), seu uso é possível. Mas, como o número de regras é $r=c^p$, este modelo torna-se inviável para um aumento de c ou p e rapidamente esgota-se a memória de um computador comum durante o treinamento. Para 6 variáveis, 3 “clusters” por variável, considerando-se 8 bytes de memória por valor armazenado e 1000 pontos coletados tem-se, apenas para a matriz Φ :

$$\begin{aligned}
 3^6 \times (6+1) &= 5.103 && \text{parâmetros} \\
 5.103 \times 1000 &= 5.103.000 && \text{valores em } \Phi \\
 5.103.000 \times 8 &= 40.824.000 && \text{bytes} \cong 40\text{MBytes}
 \end{aligned}$$

Além da dimensionalidade, outro problema ou característica do algoritmo anterior é ilustrado na Figura 16.

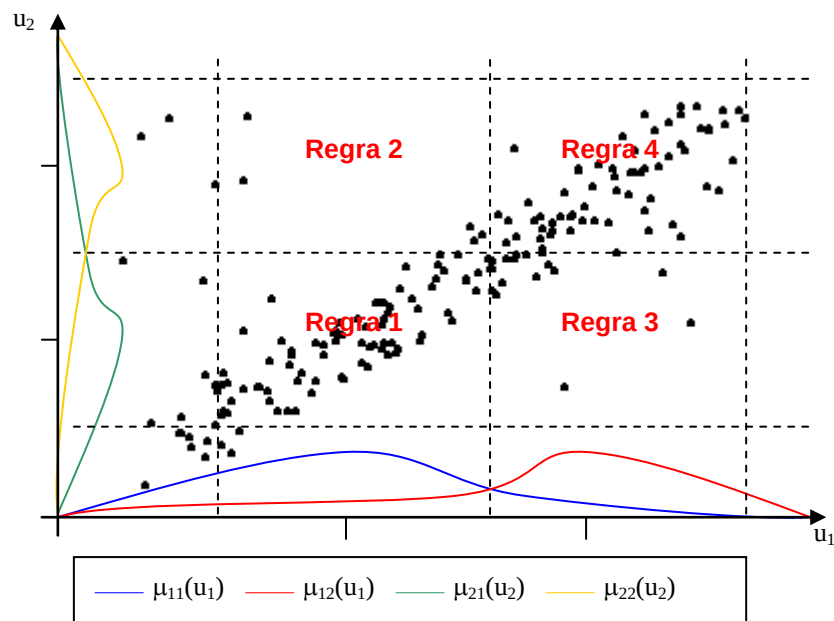


Figura 16 - Exemplo do espaço de entradas para o modelo de Wang-Langari

Pode ser observado neste exemplo, que, embora a clusterização tenha apresentado resultados corretos, só há dados para definir duas das regras existentes. Apenas as regras 1 e 4 seriam suficientes; além disso, as regras adicionais consomem tempo de processamento e memória.

Por estes motivos, é proposto um novo método/modelo a ser usado, denominado de “Regras Limitadas”. Para gerar este novo modelo, os dados coletados são agrupados como pontos num espaço de dimensão $p+1$, onde p é o número de variáveis de entrada, e então são clusterizados. Este modo de clusterização, onde todo o espaço de regressão é utilizado, é conhecido como *Product Space Clustering*.

O uso de todo o espaço de regressão na clusterização apresenta melhores resultados, pois, como a forma da função afeta a clusterização, é possível identificar melhor as regiões não lineares, atribuindo-lhes um maior número de agrupamentos. A Figura 17 ilustra esta afirmação.

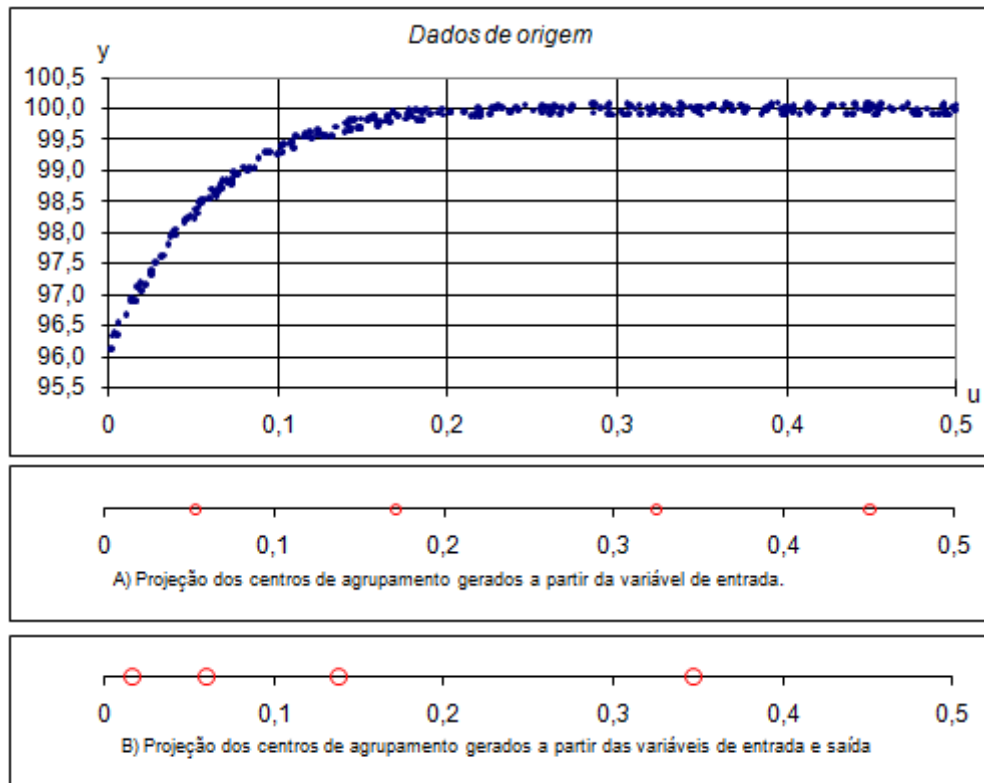


Figura 17 - Diferenças do resultado da clusterização. A primeira (A) utiliza apenas a variável de entrada, a segunda (B) utiliza ambas as variáveis (OLIVEIRA; GARCIA, 2003).

É possível observar claramente neste exemplo as diferenças entre os dois modos de clusterização. Os dados são gerados a partir de uma função e são adicionados alguns erros aleatórios. Este conjunto de dados é clusterizado pelo FCM de modo a encontrar quatro agrupamentos pelas duas formas.

Na clusterização **A** são utilizados apenas os dados relativos à variável de entrada u . Como os valores da variável u estão uniformemente distribuídos no intervalo, os centros dos agrupamentos resultaram também uniformemente distribuídos. A consequência disso é que existem dois agrupamentos para descrever a região de saída constante, $u > 0,2$, e apenas dois outros agrupamentos descrevendo a região mais complexa, $u < 0,2$.

Na clusterização **B** é utilizada a clusterização do espaço de regressão completo, ou seja, são utilizados os dados de entrada e saída juntos. É possível observar que deste modo utiliza-se apenas um agrupamento para descrever a região de saída constante, $u > 0,2$. E assim é possível utilizar a maioria dos agrupamentos para descrever a região mais complexa. Um agrupamento é designado a descrever a região

aproximadamente linear, $u < 0,04$, e os outros dois estão associados à região não linear, $0,04 < u < 0,2$. Isto ocorre, pois, nas regiões não lineares, os pontos estão relativamente mais próximos uns dos outros do que nas regiões que podem ser aproximadas por retas. Desta forma, este tipo de clusterização proporciona uma melhora na descrição do sistema, pois leva em consideração tanto a distribuição dos dados no espaço de entrada quanto o formato da função.

A desvantagem deste método é a perda do caráter linguístico da inferência nebulosa. No modelo original, do tipo Takagi-Sugeno gerado pelo método de Wang-Langari, eram encontrados agrupamentos para cada variável, o que ainda permitia a associação dos agrupamentos a valores linguísticos como “alto”, “médio” e “baixo”. Neste modelo, a princípio, não há como fazer tal associação; no entanto, isto não parece uma perda tão grande, tendo-se em vista os objetivos almejados.

Neste método, a cada agrupamento gerado é associada uma regra; dessa forma o número de agrupamentos define o número de regras diretamente, e por isso é chamado de “Regras Limitadas”. A clusterização é feita utilizando-se o FCM e, para a estimação dos parâmetros, utiliza-se o método dos mínimos quadrados.

Como no método anterior, os dados de entrada são agrupados numa matriz de entrada $\mathbf{Z}$:

$$\mathbf{Z} = \begin{bmatrix} u_1^1 & \cdots & u_p^1 & y^1 \\ \vdots & \ddots & \vdots & \vdots \\ u_1^n & \cdots & u_p^n & y^n \end{bmatrix}$$

onde $u_i^k = u_i(k)$ e $y^k = y(k)$. Mas, neste modelo, os dados de entrada são passados ao FCM como pontos pertencentes ao $\mathbf{R}^{p+1}$. Definido o número de “clusters” c , o FCM encontra c centros de “clusters” de dimensão $p+1$. Estes centros podem ser agrupados numa matriz $\mathbf{V}$, agora na forma:

$$\mathbf{V} = \begin{bmatrix} x_{11} & \cdots & x_{1(p+1)} \\ \vdots & \ddots & \vdots \\ x_{c1} & \cdots & x_{c(p+1)} \end{bmatrix}$$

onde cada linha representa as coordenadas $\mathbf{v}_i$ de um centro de “cluster”. Como durante a execução não se tem acesso à variável de saída, cujo valor deve ser estimado, convém armazenar apenas as coordenadas referentes às variáveis de entrada:

$$\mathbf{V} = \begin{bmatrix} x_{11} & \cdots & x_{1p} \\ \vdots & \ddots & \vdots \\ x_{c1} & \cdots & x_{cp} \end{bmatrix}$$

A pertinência de um ponto a um agrupamento pode ser calculada por:

$$\mu_{ik} = \frac{1}{\sum_{j=1}^c (d_{ki} / d_{kj})^{2/(m-1)}} \quad (30)$$

onde $d_{ki} = \|\mathbf{z}_k - \mathbf{v}_i\|$ é a distância do ponto k ao centro do agrupamento, $\mathbf{v}_i$. Neste modelo, a cada agrupamento é associada uma regra. Assim são geradas $r=c$ regras, cujos graus de ativação são $g_h^k = \mu_{hk}$.

A saída deste modelo nebuloso é:

$$y_f^k = \sum_{h=1}^r g_h^k y_h^k = \sum_{h=1}^r g_h^k \cdot (b_{h0} + b_{h1}u_1^k + \dots + b_{hp}u_p^k) \quad (31)$$

Para finalizar o modelo, falta apenas calcular os parâmetros consequentes b_{ij} . O objetivo é minimizar a somatória dos erros quadráticos para o modelo alimentado com os pontos coletados. Ou seja, deseja-se minimizar:

$$J = \sum_{k=1}^n (y^k - y_f^k)^2 = \sum_{k=1}^n \left(y^k - \sum_{h=1}^r g_h^k \cdot (b_{h0} + b_{h1}u_1^k + \dots + b_{hp}u_p^k) \right)^2 \quad (32)$$

Definindo-se:

$$\mathbf{Y} = [y^1 \ y^2 \ \dots \ y^n]^T$$

$$\mathbf{B} = [b_{10} \ \dots \ b_{1p} \ b_{20} \ \dots \ b_{2p} \ \dots \ b_{r0} \ \dots \ b_{rp}]^T$$

$$\Phi = \begin{bmatrix} g_1^1 & \dots & g_r^1 & g_1^1 u_1^1 & \dots & g_r^1 u_1^1 & \dots & g_1^1 u_p^1 & \dots & g_r^1 u_p^1 \\ g_1^2 & \dots & g_r^2 & g_1^2 u_1^2 & \dots & g_r^2 u_1^2 & \dots & g_1^2 u_p^2 & \dots & g_r^2 u_p^2 \\ \vdots & & & \vdots & & \vdots & & \vdots & & \vdots \\ g_1^n & \dots & g_r^n & g_1^n u_1^n & \dots & g_r^n u_1^n & \dots & g_1^n u_p^n & \dots & g_r^n u_p^n \end{bmatrix} \quad (33)$$

A função objetivo J , a ser minimizada, pode então ser reescrita como:

$$J = \|\mathbf{Y} - \Phi \cdot \mathbf{B}\|^2$$

Desta forma, chega-se ao mesmo problema do item 3.6, onde a solução foi descrita. Para um novo vetor de entradas, é necessário calcular as novas distâncias e pertinências aos agrupamentos pela equação (30). Então, determinam-se os graus de ativação de cada regra e, a partir da equação (31), estima-se a nova saída y_f .

A Figura 18 ilustra como seria dividido o espaço de entradas para os mesmos dados da Figura 16.

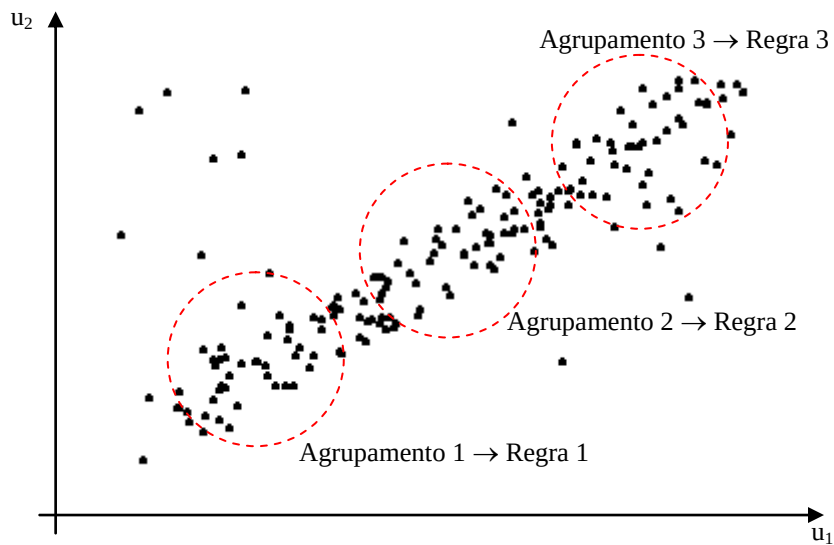


Figura 18 - Ilustração do espaço de entrada para o modelo “Regras Limitadas”.

3.7.3 Modelo gerado a partir da clusterização de Gustafson-Kessel.

O método de obtenção do modelo nebuloso, a partir da clusterização de Gustafson-Kessel, é muito parecido com o método “Regras Limitadas”. Obtém-se uma matriz de entrada $\mathbf{Z}$ conforme ocorre no item anterior, os dados são clusterizados pelo algoritmo de Gustafson-Kessel e obtêm-se os centros dos “clusters” de dimensão $p+1$, assim como as matrizes de covariância $\mathbf{F}_i$. Ignora-se a coordenada referente à variável de saída dos centros dos agrupamentos e as últimas coluna e linha da matriz $\mathbf{F}_i$, também referentes a esta variável. As distâncias d_{ik} são calculadas pela equação (18). O restante do método de obtenção do modelo e estimação da saída é igual ao “Regras Limitadas”.

Para se estimar o valor de saída para uma nova entrada, deve-se calcular os novos valores de distância e pertinência pelas equações (16) e (18). E então a saída é calculada pela equação (31).

3.8 Métodos de identificação de sistemas variantes no tempo

O principal interesse deste trabalho é criar um sensor virtual que possa estimar variáveis em sistemas dinâmicos, cujas propriedades possam mudar com o tempo, mudanças que podem ocorrer devido a várias razões.

Por isso, muitos sistemas dinâmicos requerem a disponibilidade de um modelo em tempo real conforme o processo se desenvolve. Aplicações típicas incluem controle adaptativo, predição adaptativa, detecção de falha e monitoramento do processo, tratando-se assim de processos complexos que podem ser modelados como sistemas nebulosos com algoritmos do tipo Takagi-Sugeno, descritos no item (3.3.2), que pode expressar uma relação funcional altamente não linear usando um número pequeno de regras.

Uma vez que a estrutura da premissa no modelo tipo Takagi-Sugeno, como visto nos itens (3.4) e (3.5), tenha sido determinada, a tarefa seguinte é estimar os parâmetros dos modelos de regressão lineares na parte consequente das regras, como descrito nos itens (3.6) e (3.7).

Um dos mais populares algoritmos de identificação de sistemas em tempo real é o algoritmo dos mínimos quadrados recursivos (*Recursive Least Square* - RLS). Este algoritmo é conhecido por ter ótimas propriedades quando os parâmetros do sistema são invariantes no tempo. Contudo, ele é inadequado para rastrear parâmetros variantes no tempo, uma vez que o ganho do algoritmo converge para zero, significando que depois de algumas iterações o erro de predição não mais desempenhará seu papel. Deste modo, considerável esforço de pesquisa tem sido dirigido em direção ao desenvolvimento de versões modificadas do algoritmo que evite que o ganho vá para zero prematuramente.

Porém, o algoritmo RLS pode ser modificado para tratar parâmetros variantes no tempo, introduzindo-se o chamado fator de esquecimento (*forgetting factor*), também conhecido como constante de ponderação de dados exponencial.

3.8.1 Mínimos quadrados recursivos com fator de esquecimento constante

O mais conhecido destes algoritmos RLS, modificado para poder ser usado com sistemas variantes no tempo, é aquele baseado em ponderação de dados exponenciais descrito por Goodwin e Paine (apud WANG; LANGARI, 1994).

O algoritmo RLS com esquecimento exponencial é dado por

$$\hat{\theta}_k = \hat{\theta}_{k-1} + K_k e_k \quad (34a)$$

$$e_k = y_k - \phi_k^T \hat{\theta}_{k-1} \quad (34b)$$

$$K_k = \frac{P_{k-1} \phi_k}{\lambda_k + \phi_k^T P_{k-1} \phi_k} \quad (34c)$$

$$P_k = \frac{1}{\lambda_k} (I - K_k \phi_k^T) P_{k-1} \quad (34d)$$

onde $\phi_k, \hat{\theta}_k \in \mathbb{R}^n$; $P_k \in \mathbb{R}^{n \times n}$ é a matriz de covariância, $K_k \in \mathbb{R}^n$ é o chamado ganho de Kalman; $y \in \mathbb{R}$, e λ_k é o fator de esquecimento com valor constante. Este algoritmo tem muitas propriedades de convergência desejadas, mas infelizmente pode encontrar

o problema de explosão da matriz de covariância, ou seja, a matriz P_k pode aumentar quando k cresce se o sinal de entrada é tal que a magnitude de $P_{k-1}\phi_k$ seja pequena (se o sinal de entrada não é persistentemente excitante) (AGUIRRE, 2004) (WANG; LANGARI, 1996b).

3.8.2 Mínimos quadrados recursivos com fator de esquecimento variável

Muitos algoritmos, modificados, foram propostos para superar o problema de explosão da matriz de covariância, dentre os quais o algoritmo de Fortescue et al. (WANG; LANGARI, 1996b) é o mais interessante e encorajador, apesar de apresentar algumas desvantagens.

3.8.2.1 Algoritmo de Fortescue

No algoritmo de Fortescue et al. - FKY, o fator de esquecimento λ_k , observado na equação (34c) e (34d), é variante no tempo e seu valor é determinado pela seguinte regra de ajuste:

$$\lambda_k = 1 - (1 - \phi_k^T K_k) e_k^2 / \Sigma_0 \quad (35)$$

onde Σ_0 é uma constante de ajuste positiva.

Claramente, este algoritmo ajusta o valor do fator de esquecimento baseado no monitoramento do quadrado do erro de predição, que é tido como um dos mais convenientes valores para a supervisão do desempenho de um algoritmo de identificação (ISERMAN; LACHMAN, 1985 apud WANG; LANGARI, 1996b). Quando o quadrado do erro de predição aumenta, λ_k é reduzido.

De qualquer modo, uma possível desvantagem deste algoritmo é que ele pode dar resultados muito conservadores. A razão é que quando os parâmetros do sistema são alterados, o fator de esquecimento do algoritmo convergirá para a unidade em um

período de tempo muito pequeno, ou seja, o algoritmo é reduzido a um RLS comum sem fator de esquecimento; por isso ele não pode proporcionar esquecimento exponencial contínuo, ou, ainda, os dados antigos podem não ser suficientemente esquecidos.

3.8.2.2 Algoritmo de Wang-Langari

Wang e Langari (WANG; LANGARI, 1996b) propuseram uma modificação no algoritmo de Fortescue et al. – FKY, produzindo uma nova regra de ajuste do fator de esquecimento. A diferença essencial entre este novo algoritmo e o algoritmo FKY é que o fator de esquecimento convergirá para uma constante menor que a unidade no regime permanente, em vez da constante um. Deste modo, as informações baseadas em medidas passadas são sempre parcialmente esquecidas, ao passo que às informações baseadas em novas medidas é sempre dada considerável atenção.

A modificação introduzida no algoritmo de FKY é simplesmente a seguinte:

$$\lambda_k = \Delta - \Gamma e_k^2 \quad (36)$$

com

$$0 < \Delta < 1, \quad \Gamma > 0$$

onde Δ e Γ são dois fatores de ajuste com valores constantes. Na prática, o fator Δ é planejado para ser menor que mas próximo à unidade, para proporcionar esquecimento exponencial contínuo, enquanto isso evita o problema de explosão da covariância. O fator Γ é projetado para ser um número grande, sendo o escolhido 1000. As numerosas simulações, feitas por Wang e Langari, mostraram que a escolha de Γ é muito menos sensível que a escolha de Σ_0 no algoritmo FKY.

Wang e Langari (WANG; LANGARI, 1996b) provaram que no regime permanente o algoritmo é reduzido ao algoritmo RLS ordinário com fator de esquecimento constante, isto é, o algoritmo mostrado nas equações (34a) a (34d). Poderia ser argumentado ainda que este algoritmo enfrenta o mesmo problema de explosão da covariância, mas felizmente o valor do fator de esquecimento em regime

permanente ($\lambda_k = \Delta$) pode ser planejado para ter um valor grande o suficiente (ou seja, 0,995) comparado com aqueles do algoritmo RLS ordinário com fator de esquecimento constante; assim o problema de explosão da covariância raramente aparecerá.

Assim, o algoritmo combina as vantagens do algoritmo FKY e o algoritmo RLS ordinário com fator de esquecimento constante; quando o erro de predição aumenta inesperadamente, λ_k rapidamente diminuirá para se adaptar a mudanças, produzindo um grande esquecimento exponencial; quando o erro de predição tende a zero, λ_k se aproxima da constante Δ , para garantir um moderado, mas contínuo, esquecimento exponencial.

A fim de evitar que o fator de esquecimento seja muito grande ou muito pequeno, até mesmo negativo, uma restrição quanto ao limite superior e ao limite inferior é acrescentada:

$$\lambda_k = \begin{cases} \lambda_{\min}, & \text{se } \lambda_k < \lambda_{\min} \\ \lambda_{\max}, & \text{se } \lambda_k > \lambda_{\max} \end{cases} \quad (37)$$

Comumente, $\lambda_{\max}$ é ajustado como 1 (unidade), e $\lambda_{\min}$ é ajustado como um número pequeno, como por exemplo 0,2.

A superioridade adicional deste algoritmo sobre o algoritmo RLS ordinário é essencialmente um peso adicional de atualização em cada passo de tempo, consistindo de uma operação escalar simples. Tal aumento na complexidade computacional é insignificante, comparado com aquela do algoritmo FKY.

O algoritmo de Wang e Langari em uma forma mais compacta é:

$$\hat{\theta}_k = \hat{\theta}_{k-1} + \frac{P_{k-1}\phi_k e_k}{\lambda_k + \phi_k^T P_{k-1} \phi_k} \quad (38a)$$

$$e_k = y_k - \phi_k^T \hat{\theta}_{k-1} \quad (38b)$$

$$P_k = \frac{1}{\lambda_k} \left(I - \frac{P_{k-1}\phi_k\phi_k^T}{\lambda_k + \phi_k^T P_{k-1} \phi_k} \right) P_{k-1} \quad (38c)$$

$$\lambda_k = \Delta - \Gamma e_k^2, \quad \lambda_{\min} \leq \lambda_k \leq \lambda_{\max} \quad (38d)$$

3.9 Equacionamento utilizado na implementação numérica

O algoritmo modificado parece ser bastante eficiente em identificação de sistemas variantes no tempo, mas ele pode encontrar dificuldade numérica durante sua implementação prática. O principal problema do algoritmo é que a subtração de matrizes em (38c), que representa a redução em incertezas devido à medida, pode produzir uma matriz P_k que é computacionalmente não positiva definida, uma impossibilidade teórica. Para superar esta dificuldade, várias técnicas numericamente confiáveis, tais como a decomposição em raiz quadrada, a fatorização UD e a decomposição em valor singular, podem ser usadas para o cálculo de P_k . Será apresentado aqui o método computacionalmente simples de decomposição em raiz quadrada que é usado para reformular o algoritmo. C_k , uma matriz triangular inferior, denota a decomposição em raiz quadrada de P_k , isto é

$$P_k = C_k C_k^T \quad (39)$$

Então o algoritmo de (38a) – (38d) pode ser reformulado e será apresentado a seguir (WANG; LANGARI, 1996b):

$$\begin{aligned} \mathbf{f}_k &= C_{k-1}^T \phi_k \\ \gamma_k &= [\lambda_k + \mathbf{f}_k^T \mathbf{f}_k]^{1/2} \\ \alpha_k &= \gamma_k + \lambda_k^{1/2} \\ \beta_k &= 1/(\gamma_k \alpha_k) \\ \sigma_k &= \beta_k C_{k-1} \mathbf{f} \\ \mathbf{e}_k &= y_k - \phi_k^T \hat{\theta}_{k-1} \\ \hat{\theta}_k &= \hat{\theta}_{k-1} + \sigma_k (\alpha_k / \gamma_k) \mathbf{e}_k \\ \lambda_k &= \Delta - \Gamma \mathbf{e}_k^2 \\ \text{ajustar } \lambda_k &= \begin{cases} \lambda_{\min}, & \text{se } \lambda_k < \lambda_{\min} \\ \lambda_{\max}, & \text{se } \lambda_k > \lambda_{\max} \end{cases} \\ C_k &= \frac{1}{\lambda_k^{1/2}} (C_{k-1} - \sigma_k \mathbf{f}_k^T) \end{aligned} \quad (40)$$

4 PROCESSO DE NEUTRALIZAÇÃO DE pH

4.1 Descrição de um processo de neutralização de pH

É possível encontrar diversas aplicações onde é observada a importância da medição, regulação e controle do pH. Entre elas podem-se citar plantas de tratamento de água potável, processamentos biológicos, processamento químico, fabricação de papel, processamento do vidro, tratamento de efluentes residuais, e também as mostradas na Tabela 2, que comprovam esta consideração (SOTOMAYOR, 1997).

Tabela 2 - Aplicação do pH em produções industriais

Aplicação	Processos e aspectos afetados pelo pH
Bacteriológica	Cultivo de microorganismos e metabolismo
Cozimento	Volume de massa, textura e cor
Cervejaria	Produção de extrato e açúcar durante o <i>mashing</i>
Conservas	Tempo e temperatura de esterilização
Química	Impurezas e cristalização de sais
Limpeza	Remoção de pinturas
Tintas	Produção e uniformidade de cores
Eletroblindagem	Niquelamento de depósitos
Fermentação	Tempo de fermentação e cultivo de organismos
Gelatina	Absorção de água, solubilidade e claridade
Farmacêutica	Estabilidade e reação do corpo humano
Pigmentos	Uniformidade da composição em precipitação
Águas residuais	Tempo de formação, cheiro e <i>foaming</i>
Têxtil	Processos úmidos
Tratamento de água	Coagulação e processos de abrandamento

Um processo de neutralização de pH pode ser funcionalmente representado pela Figura 19.

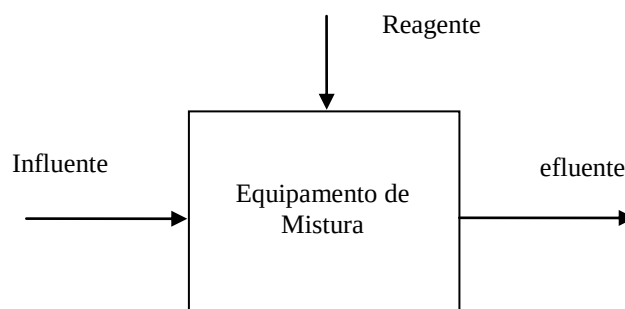


Figura 19 - Uma ilustração de um processo de neutralização de pH

O fluído de entrada cujo pH será ajustado (neutralizado dentro do misturador) é chamado de influente; é chamado de reagente o ácido (ou base) usado para ajustar o pH, e de efluente, a saída com o pH ajustado.

Um sistema de controle de pH típico composto de eletrodos de pH, transmissor de pH, um controlador, válvula de controle e um equipamento de mistura é funcionalmente representado na Figura 20, adaptado de Sotomayor (1997).

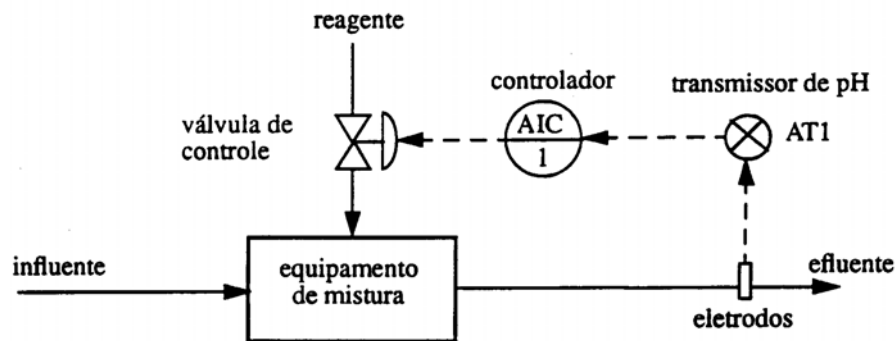


Figura 20 - Sistema de controle típico de pH

Embora um processo de pH real tenda a ser muito complicado devido às variações das espécies contidas no influente e no reagente e da seleção do equipamento de mistura, existem vários modelos matemáticos dinâmicos conhecidos representando as características dominantes de um processo de pH em uma unidade CSTR (*continuous stirred tank reactor*).

Estes modelos foram desenvolvidos com base em princípios fundamentais como balanço de massa e equilíbrio químico e foram experimentalmente comprovados. Portanto, a planta de neutralização de pH que será utilizada nesta investigação é uma unidade CSTR. Esta planta tem como objetivo a obtenção de um efluente com pH desejado (na maioria das aplicações, neutro) através da titulação de uma base (fluido titulante ou neutralizante) em um ácido (fluido genérico a ser tratado).

Serão apresentadas a seguir duas possíveis representações matemáticas (modelos) que já foram e continuam sendo bastante utilizadas e testadas na literatura científica, e também a justificativa para utilização de uma delas no desenvolvimento deste trabalho.

4.1.1 Modelo com apenas uma variável de saída (pH)

McAvoy (apud NIE; LOH; HANG, 1996) (PAIVA; MANZI, 2007) desenvolveu um modelo dinâmico que foi verificado experimentalmente, onde é assumido que na unidade CSTR ocorre uma mistura perfeita e isotérmica.

Este modelo consiste de duas partes: um modelo dinâmico descrevendo o fluxo (dinâmico) de concentrações das composições químicas na unidade CSTR; um modelo estático não linear caracterizando as condições de equilíbrio físico-químico entre estas concentrações.

Assumindo que o processo de neutralização de pH possua dois fluidos entrando na unidade CSTR, sendo que o fluxo de entrada F_a contém um ácido com concentração C_a e o outro fluxo F_b contém uma base com concentração C_b , tem-se o seguinte modelo dinâmico para a unidade CSTR:

$$V \frac{dw_a}{dt} = F_a C_a - (F_a + F_b) w_a \quad (41a)$$

$$V \frac{dw_b}{dt} = F_b C_b - (F_a + F_b) w_b \quad (41b)$$

Onde a constante V é o volume da mistura contida no reator, e w_a e w_b são respectivamente as concentrações de ácidos e bases não reagentes contidas no tanque. As equações (41a) e (41b) descrevem como as concentrações w_a e w_b mudam dinamicamente com o tempo, sujeitas aos fluxos F_a e F_b .

Para obter o pH no efluente, é necessário encontrar uma relação entre as concentrações instantâneas w_a e w_b e valores do pH. Esta relação pode ser descrita por uma equação algébrica não linear conhecida como curva de titulação ou curva característica e dependendo das espécies químicas usadas, a curva de titulação varia.

Como exemplo, consideremos dois casos:

- (a) Um influente fraco é neutralizado por um reagente forte;
- (b) Um influente forte é neutralizado por um reagente forte.

Considerando o ácido acético (ácido fraco), indicado por HAC , sendo neutralizado por uma base forte $NaOH$ (hidróxido de sódio) em solução aquosa, temos as reações (NIE; LOH; HANG, 1996):



De acordo com as condições de eletroneutralidade, a soma das cargas de todos os íons na solução deve ser zero:

$$[Na^+] + [H^+] = [OH^-] + [AC^-] \quad (43)$$

Onde o símbolo $[X]$ indica a concentração do íon X .

Por outro lado, as seguintes relações de equilíbrio são válidas para a água e o ácido acético:

$$K_a = \frac{[AC^-][H^+]}{[HAC]} \quad (44a)$$

$$K_w = [H^+][OH^-] \quad (44b)$$

Onde K_a e K_w são as constantes de dissociação do ácido acético e da água com $K_a = 1,76 \times 10^{-5}$ e $K_w = 10^{-14}$. Definindo $w_a = [HAC] + [AC^-]$, $w_b = [Na^+]$ e inserindo as equações (44a) e (44b) na equação (43), obtem-se:

$$[H^+]^3 + [H^+]^2 \{K_a + w_b\} + [H^+] \{K_a(w_b - w_a) - K_w\} - K_a K_w = 0 \quad (45)$$

Usando a equação $pH = -\log_{10}[H^+]$, a equação (45) torna-se

$$w_b + 10^{-pH} - 10^{pH-14} - \frac{w_a}{1 + 10^{pK_a - pH}} = 0 \quad (46)$$

Onde $pK_a = -\log_{10} K_a$.

A relação entre w_a , w_b e pH definida pela equação (46) é referida como função de titulação e o gráfico mostrando explicitamente como o pH é determinado através da razão de w_b/w_a é referido como curva de titulação, que é mostrada na Figura 21, onde é observada claramente a não linearidade do sistema.

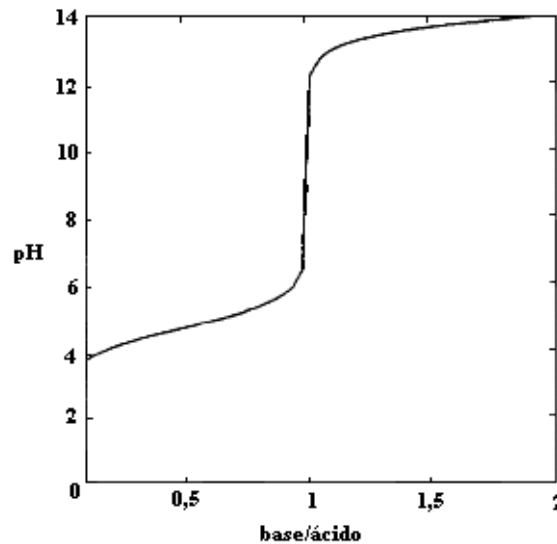


Figura 21 - Curva de titulação para ácido fraco neutralizado por base forte (NIE; LOH; HANG, 1996).

Consideremos agora o caso onde um ácido forte (influyente), por exemplo HCl , é neutralizado por uma base forte (reagente), por exemplo $NaOH$. A estequiometria pode ser descrita pelas seguintes reações (NIE; LOH; HANG, 1996):



Aplicando a condição de eletroneutralidade e a relação de equilíbrio da equação (44b) produz

$$w_b + [H^+] = \frac{K_w}{[H^+]} + w_a \quad (48)$$

Onde w_a e w_b denotam as concentrações dos íons $[Cl^-]$ e $[Na^+]$, respectivamente. A função da titulação é dada por

$$w_b - w_a + 10^{-pH} - 10^{pH-14} = 0 \quad (49)$$

E a curva de titulação correspondente é mostrada na Figura 22.

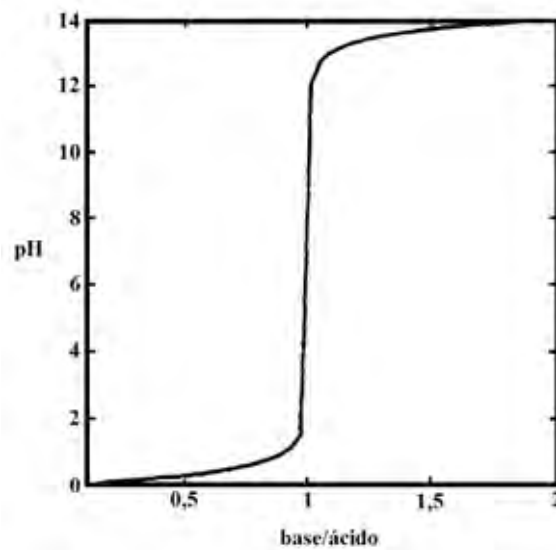


Figura 22 - Curva de titulação para ácido forte neutralizado por base forte (NIE; LOH; HANG, 1996).

Observe que a não linearidade, exibida pelo processo de neutralização de um ácido forte por uma base forte, é muito mais severa do que a de um ácido fraco neutralizado por uma base forte, como mostrado na Figura 21.

4.1.2 Modelo com duas variáveis de saída (pH e nível)

Um diagrama esquemático simplificado do processo de neutralização de pH, da unidade CSTR, é mostrado na Figura 23 (SOTOMAYOR, 1997). O processo consiste em um fluxo (q_1) de influente (ácido nítrico, HNO_3), fluxo (q_2) de “buffer” ou solução tampão (bicarbonato de sódio, $NaHCO_3$), que tem como objetivo amortecer as variações de pH, e fluxo (q_3) de titulante (mistura de hidróxido de sódio, $NaOH$, e bicarbonato de sódio, $NaHCO_3$, não mostrado na figura), os quais são misturados no

tanque TQ_4 , sendo que o nível de líquido (h_4) e o pH do efluente (pH) são as variáveis medidas.

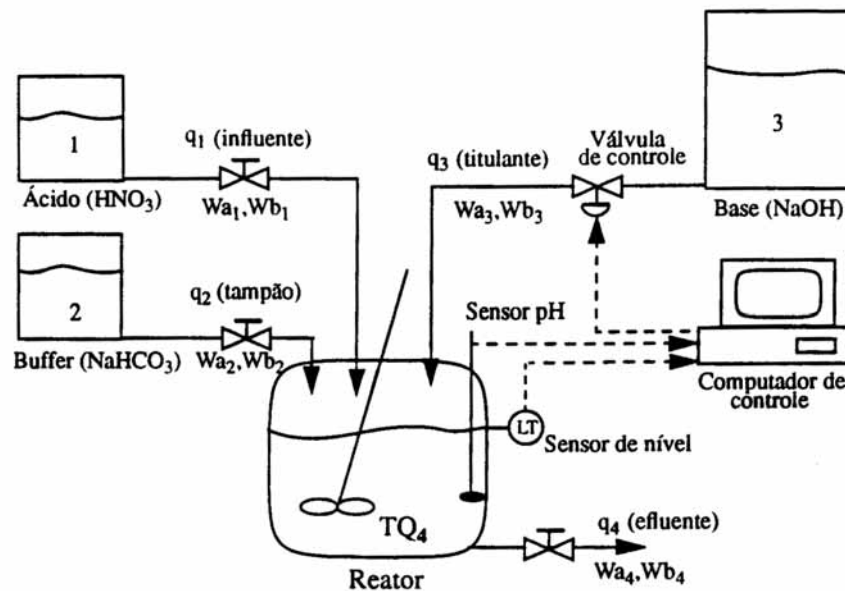


Figura 23 - Diagrama geral do processo de neutralização de pH

Uma característica que dificulta a formulação e identificação de um modelo matemático para o processo de pH é que uma pequena quantidade de elementos intrusos, como exemplo, carbono ou potássio com capacidade de tamponamento (*buffering*), mudam a dinâmica do processo.

O efeito tampão (*buffering*) é o fenômeno que ocorre se quantidades equivalentes de uma base forte e um ácido fraco, ou um ácido forte e uma base fraca, são misturados em um tanque, e a solução resultante não é neutra, ou seja, a capacidade de uma solução para resistir a mudanças no pH.

Modelos matemáticos gerais para o processo de neutralização de pH são disponíveis em uma variedade de formulações, como por exemplo, modelo de invariantes de reação, modelo de concentração de íons e o modelo com curva de titulação. Dentre estes, o modelo de invariantes de reação é considerado uma das mais completas descrições para o sistema de pH (HALL; SEBORG, 1989).

No sistema em estudo, representado pelo diagrama geral do processo de neutralização de pH, Figura 23, cada um dos fluxos de entrada contém uma certa quantidade de íons carbonato, uma base fraca. Os seguintes íons mais o ácido fraco estão contidos nos fluxos reagentes:

H^+	íon hidrogênio	OH^-	íon hidróxido
N_a^+	íon sódio	NO_3^-	íon nitrato
H_2CO_3	ácido carbônico	HCO_3^-	íon bicarbonato
CO_3^{2-}	íon carbonato		

Já que o íon sódio e o íon nitrato são conjugados de uma base forte e um ácido forte, respectivamente, isso não significa que quantidades de hidróxido de sódio e ácido nítrico sejam formadas. Consequentemente, as principais reações químicas para o sistema são (ALVAREZ et al., 2001):



As correspondentes constantes de equilíbrio são:

$$K_{a_1} = \frac{[HCO_3^-][H^+]}{[H_2CO_3]} \quad (51a)$$

$$K_{a_2} = \frac{[CO_3^{2-}][H^+]}{[HCO_3^-]} \quad (51b)$$

$$K_w = [H^+][OH^-] \quad (51c)$$

O equilíbrio químico é modelado por definição de duas variantes de reações, sendo neste caso, para cada fluxo ($i \in [1,4]$), considerado volume constante, mistura perfeita, densidade constante e completa solubilidade dos íons envolvidos. $i = 1, 2, 3$ e

4 representam as quatro vazões volumétricas a serem consideradas, respectivamente: ácido, solução tampão, base e efluente.

Portanto tem-se:

$$W_{ai} = [H^+]_i - [OH^-]_i - [HCO_3^-]_i - 2[CO_3^{2-}]_i \quad (52a)$$

$$W_{bi} = [H_2CO_3]_i + [HCO_3^-]_i + [CO_3^{2-}]_i \quad (52b)$$

A invariante W_a representa um balanço de carga relativa e a invariante W_b representa um balanço na concentração dos íons CO_3^{2-} . Ao contrário do pH, estas quantidades são preservadas. O pH pode ser determinado pela seguinte equação (HALL; SEBORG, 1989) (ALVAREZ, H. et al., 2001):

$$W_b \frac{\frac{K_{a1}}{[H^+]} + \frac{2K_{a1}K_{a2}}{[H^+]^2}}{1 + \frac{K_{a1}}{[H^+]} + \frac{K_{a1}K_{a2}}{[H^+]^2}} + W_a + \frac{K_w}{[H^+]} - [H^+] = 0 \quad (53)$$

A equação (53) representa a relação não-linear entre W_a , W_b e a concentração de íons de hidrogênio H^+ para o sistema. O pH da solução é obtido usando-se a equação (3) e repetida aqui novamente:

$$pH = -\log_{10}[H^+] \quad (54)$$

A parte dinâmica do processo de neutralização de pH é desenvolvida a seguir. Através de um balanço de massa total no reator, ou seja, no tanque (TQ_4), mostrado na Figura 23, e considerando que as massas específicas dos fluidos de entrada, de saída e do interior do reator sejam constantes e iguais, considerando também que $V = Ah$, tem-se:

$$A_4 \frac{dh}{dt} = q_1 + q_2 + q_3 - q_4 \quad (55)$$

sendo que A_4 é a área da seção transversal do reator, ou seja, tanque (TQ_4), h é a altura da solução em seu interior (nível de líquido) e q_4 é o fluxo de saída do efluente.

Dado que as invariantes de reações são por natureza preservadas, equações de conservação são usadas para representar a dinâmica de estados. Por combinação de balanços de massa em cada uma das espécies iônicas do sistema, as seguintes equações diferenciais para as invariantes de reações do efluente (W_{a4} , W_{b4}), considerando-se h variável, são derivadas (ALVAREZ et al., 2001):

$$A_4 h \frac{dW_{a4}}{dt} = q_1(W_{a1} - W_{a4}) + q_2(W_{a2} - W_{a4}) + q_3(W_{a3} - W_{a4}) \quad (56)$$

$$A_4 h \frac{dW_{b4}}{dt} = q_1(W_{b1} - W_{b4}) + q_2(W_{b2} - W_{b4}) + q_3(W_{b3} - W_{b4}) \quad (57)$$

Os sensores de pH apresentam tempos de resposta relativamente baixos quando comparados com os tempos envolvidos no processo, portanto considera-se desprezível a dinâmica no medidor de pH (SOTOMAYOR, 1997) (ALVAREZ et al., 2001). Também é desprezível a dinâmica do medidor de nível.

Para uma maior precisão na dinâmica do processo, foram introduzidos um tempo morto de medição do pH (θ_{pH}) e um atraso na mistura (θ_{mix}), expresso em (58)

$$\theta_{mix} = \frac{V}{2(F_a + \sum q_i)} \quad (58)$$

sendo que V é o volume do tanque (reator), F_a indica a taxa de descarga do agitador e $\sum q_i$ é a soma dos fluxos de entrada ($i = 1, 2, 3$) ao CSTR. F_a pode ser calculado por:

$$F_a = 7,4813 \times N_i \times N_q \times (D_i)^3 \quad (59)$$

onde:

N_i = Velocidade da hélice do agitador, dada em *rpm*.

N_q = Coeficiente de descarga do agitador; pode assumir valor entre (0,4 -0,8)

D_i = diâmetro do círculo formado pela rotação da hélice do agitador. D_i pode assumir valor entre $(0,25 - 0,4) D_t$. D_t é o diâmetro interno do CSTR.

Portanto, o tempo morto total do sistema será dado por:

$$\theta_T = \theta_{pH} + \theta_{mix} \quad (60)$$

No caso geral, considerando-se um sistema com todas as possíveis variáveis de entrada e de saída do processo de pH, tem-se:

$$u = [q_1 \quad q_2 \quad q_3 \quad q_4 \quad [C_1] \quad [C_2] \quad [C_3] \quad T_p]^T \quad (61)$$

$$y = [h \quad pH]^T \quad (62)$$

onde:

u representa as entradas do sistema a ser considerado, com as seguintes variáveis:

q_1 é o fluxo de ácido,

q_2 é o fluxo de solução tampão,

q_3 é o fluxo de base,

q_4 é o fluxo de saída,

$[C_1]$ é a concentração do fluxo de ácido - $[HNO_3]$,

$[C_2]$ é a concentração do fluxo de solução tampão - $[NaHCO_3]$,

$[C_3]$ é a concentração do fluxo de base - $[NaOH]$,

T_p é a temperatura do processo (dada em °C): T_p realmente afeta o pH, muito embora no modelo apresentado ela não tenha sido considerada,

y representa as saídas do sistema, sendo respectivamente o nível de líquido h e o **pH**.

4.2 Desenvolvimento de um modelo virtual para o processo de neutralização de pH

Uma das principais bases para o conhecimento científico e tecnológico é a construção de modelos. Segundo Ljung e Torkel (1994), o modelo é uma ferramenta usada para a obtenção de informações sobre um sistema sem a necessidade de

realização de experimento. Da mesma forma que existem diferentes tipos de sistemas, existem diferentes tipos de modelos, a saber: modelos físicos, como protótipos e plantas pilotos, modelos mentais usados para executar tarefas do cotidiano das pessoas, modelos gráficos, que são usados para descrever o comportamento do sistema por tabelas numéricas ou curvas de desempenho, e finalmente os modelos matemáticos, que podem ser definidos como uma representação abstrata da realidade por equações (CAMPOS, 2007).

Modelagem matemática e identificação constituem uma parte essencial da investigação e análise de processos do mundo real, e por isso tornaram-se objeto de atenção nas últimas décadas. Com o uso da descrição matemática de qualquer processo, é possível estudá-lo, melhorá-lo e também controlá-lo (VOUTCHKOV; VELEV, 1997).

A obtenção de modelos para sistemas de neutralização e, especificamente, para o controle de pH tem-se tornado mais relevante na biotecnologia, devido à necessidade de melhoria da qualidade do produto ou de aperfeiçoamento do processo de produção. Muito embora os fundamentos físico-químicos e a natureza eletroquímica tenham sido bem estabelecidos, o controle de tais processos ainda não está bem resolvido do ponto de vista industrial (CASILLO, 2009).

Toda vez que a experimentação num processo real apresenta restrições de ordem operacional, econômico-financeira ou de segurança, a realização de estudos de simulação a partir de um modelo do processo é fundamental, seja com o objetivo de treinamento, projeto ou predição de resultados (SANTOS, 2007).

O processo de neutralização de pH, uma etapa comum em muitos processos químicos como parte de tratamentos de águas residuais, tem sido amplamente considerada como um problema muito difícil em controle de processos. A complexidade deste processo é maior quando, ao invés de um único ácido, o efluente apresenta uma combinação de ácidos (uma situação mais realista). Vários exemplos de sucesso de novas estratégias de controle foram apresentados na literatura, mas o problema não foi ainda totalmente resolvido.

Algumas propostas foram testadas apenas em um processo simulado, enquanto outras foram testadas num banco de ensaios experimentais, tanto em escala de laboratório como em planta piloto.

Contudo, um importante problema resulta da falta de congruência e homogeneidade no que diz respeito às condições de operação, restrições, perturbações aplicadas, e à capacidade de reproduzir os processos industriais.

Uma alternativa para superar este problema é determinar um modelo de referência (*benchmark model*) que permita a realização de estudos comparativos de diferentes estratégias de controle. Modelos de referência são regularmente utilizados em outras áreas científicas, embora na área de controle de processo apenas alguns poucos modelos de referência estejam disponíveis (BUPP; BERNSTEIN; COPPOLA, 1995).

Em particular, apenas o artigo de Alvarez et al. (2001) mostra um modelo de referência para controle de neutralização do pH, ainda que muitos trabalhos tenham sido relatados neste campo.

É importante notar que um modelo de referência é mais que um modelo matemático ou uma montagem de laboratório. Um modelo de referência para projeto de controle não-linear inclui um modelo matemático, uma montagem de laboratório, um conjunto de condições para o funcionamento do processo, um conjunto de requisitos explícitos para o processo controlado, um índice para a avaliação do desempenho do controlador, e um conjunto de sinais de teste a ser aplicado durante os ensaios. Além disso, o modelo matemático deve ser validado com a montagem de laboratório. Uma condição desejável é que o equipamento de laboratório também seja validado com uma planta industrial real. No artigo de Alvarez et. al. (2001), essas condições foram satisfeitas, médias de medições do processo de neutralização de plantas industriais são tomadas e assim dimensionada para um processo condizente à um equipamento de laboratório, e, finalmente, tal montagem é validada em relação a um modelo matemático proposto. Todas estas bases de referência indicadas são tomadas diretamente de trabalhos anteriores (HALL; SEBORG, 1989) (HENSON; SEBORG, 1994), com a aplicação das modificações necessárias para apresentar um modelo de referência formal para o processo de neutralização de pH real. Tais

modificações foram aplicadas após uma extensa revisão de plantas de neutralização instaladas em vários países. Apesar de suas condições de funcionamento serem diferentes daquelas mencionadas em Hall e Seborg (1989), os mesmos serviram como base para formular um modelo de referência. Um modelo de referência bem determinado pode ser definido como simples, porque seu modelo matemático e procedimento de operação são simples, embora o mesmo incorpore todas as características importantes do processo; e digno de crédito, uma vez que reproduz a maioria das dificuldades reais de um dado problema industrial; é também independente, por conter informações suficientes do processo que permitem simulação do modelo sem cálculos extensos; e claro, com os mecanismos do processo coerentes e de fácil interpretação física, permitindo que os resultados numéricos da simulação sejam analisados em relação às condições reais do processo. Portanto, foi desenvolvido em Simulink[®]/Matlab[®] um modelo virtual para o processo de neutralização de pH em um tanque CSTR, baseado nos trabalhos descritos anteriormente, que representa um modelo de referência na literatura técnico-científica, com reconhecida utilidade e aplicação. O apêndice A contém a descrição em Matlab[®]/Simulink[®] para o processo.

4.2.1 Intervalos de operação e indicação do modelo

Com o objetivo de testar as principais características dos sensores virtuais apresentados nos capítulos anteriores, foram realizados exemplos de simulação utilizando o ambiente de programação Matlab[®]/Simulink[®].

Com base no equacionamento desenvolvido no item 4.1.2, e com o propósito de criar um modelo virtual, ou seja, um modelo computacional para o modelo de referência descrito, é importante definir um conjunto de intervalos de operação que permita ao modelo reproduzir condições encontradas em processos de neutralização de pH industriais típicos, mas em uma bancada em escala de laboratório. Valores médios, que se aproximem dos valores da planta verdadeira, são necessários. Para conseguir tal feito, foi necessária a colaboração de experientes engenheiros de processo com pesquisadores, que resultou na montagem de uma bancada de teste do processo em

escala de laboratório. A Tabela 3 mostra as condições nominais de estado estável (*nominal stable-state* – NSS) para pH = 7,0 em um processo industrial de referência e em uma bancada de teste em escala de laboratório (ALVAREZ et al., 2001).

Tabela 3 – Instalação industrial e bancada de laboratório sob as condições normais de estado estável NSS

Parâmetros	Indústria	Laboratório
Volume V	14 m ³	0,02 m ³
diâmetro	2,612 m	0,294 m
Ácido	mistura	Forte (HNO₃)
K _{a1}	---	4,47 x 10 ⁻⁷
K _{a2}	---	5,62 x 10 ⁻¹¹
K _w	1,0 x 10 ⁻¹⁴	1,0 x 10 ⁻¹⁴
q _{1_NSS}	0,01166 m ³ /s	1,666 x 10 ⁻⁵ m ³ /s
T _{RESA}	1200 s	1200 s
W _{a1} =[HNO ₃]	---	0,01 M
W _{b1} =cte	---	0 M
Buffer-tampão	---	NaHCO₃
q _{2_NSS}	---	5,5 x 10 ⁻⁷ m ³ /s
[NaHCO ₃]	---	0,1 M
W _{a2} =cte	---	-0,1 M
W _{b2} =cte	---	0,1 M
base	NaOH	NaOH
q _{3_NSS}	5,83 x 10 ⁻³ m ³ /s	7,81 x 10 ⁻⁶ m ³ /s
Composição molar de base	~0,025 NaOH e 3,5 x 10 ⁻⁴ NaHCO ₃	~0,025 NaOH e 3,5 x 10 ⁻⁴ NaHCO ₃
W _{a3} =cte	---	-0,02535 M
W _{b3} =cte	---	0,00035 M
q _{3CMAX}	---	2,4 x 10 ⁻⁵ m ³ /s
q _{4_NSS}	~1050 x 10 ⁻⁵ m ³ /s	2,5 x 10 ⁻⁵ m ³ /s
W _{a4_NSS}	---	-0,0018862 M
W _{b4_NSS}	---	0,0023067 M
L	2.0 m	0,446 m
D _p	0,15 m (~5")	0,012 m (~0,5")
θ _{TTE}	2 s	2 s
C _{MIX}	0,05	0,01
θ _{MIX}	40 s	8 s
θ _{pH}	42 s	10 s
T _{EPC}	10 s	10 s

Com o conjunto de dados acima e as equações 53 e 56-60 é possível acompanhar a evolução temporal do pH do fluxo de saída, ou seja, o pH do efluente.

Com base nas discussões anteriores, e com os resultados relatados em artigos (HALL; SEBORG, 1989) (HENSON; SEBORG, 1994) (ALVAREZ et al., 2001) e

livros (SEBORG; EDGAR; MELLICHAMP, 1989) (SHINSKEY, 1996) sobre processos de neutralização de pH, pode-se afirmar que a neutralização de um ácido forte (influyente) por uma base forte (titulante) representa um problema multivariável difícil, porque cada variável manipulada tem um efeito significativo em cada saída controlada. Assim, esses processos são altamente não lineares e variantes no tempo, devido a não linearidade inerente associada ao controle de pH e às mudanças na curva de titulação, que ocorrem quando a quantidade de agentes de *buffer* se alteram de maneira imprevisível.

A estabilidade do sistema deve ser garantida para a operação em malha fechada sob as condições particulares e perturbações indicadas pelo modelo de referência. O modelo de referência determina duas classes de perturbações, degraus variáveis e ondas senoidais periódicas, que se aproximam do que ocorre em sistemas industriais onde uma tubulação principal recebe resíduos de vários processos. As condições específicas para a escolha de tais perturbações são as seguintes:

Degrau variável. Esta condição é geralmente encontrada quando a tubulação de efluente é curta ou tem um diâmetro pequeno ou não há um tanque que suavize o fluxo antes do tanque de neutralização. Apesar das mudanças bruscas da variável, este tipo de perturbação é frequentemente encontrada em instalações industriais que possuem instalações e equipamentos de baixo custo.

Ondas senoidais periódicas. Esta perturbação surge quando um tanque de suavização de fluxo é colocado antes do tanque de neutralização, ou quando a tubulação do efluente tem o diâmetro e comprimento suficientes para funcionar como um reservatório para suavização do fluxo. Em ambos os casos, a típica homogeneização e os efeitos de diluição estão bem representados por ondas senoidais de período e amplitude constantes.

A seguir, serão apresentados os resultados das simulações propostas nos artigos de Hall (HALL; SEBORG, 1989) e Alvarez (ALVAREZ et al., 2001) para as duas perturbações descritas.

As respostas de malha aberta do processo para os testes 1, Tabela 4, e teste 2, Tabela 5, são apresentadas nas Figuras 24 e 25, respectivamente.

Tabela 4 – Teste 1: degraus de q_1 com diferentes valores de *buffer*

tempo (min)	q_1 (l/min)	W_{a1} (M)	q_2 (l/min)
0	1,3 q_{1_NSS}	0,01	q_{2_NSS}
40	0,7 q_{1_NSS}	0,01	q_{2_NSS}
80	1,3 q_{1_NSS}	0,01	0
120	0,7 q_{1_NSS}	0,01	0
160	q_{1_NSS}	0,01	4 q_{2_NSS}
200	q_{1_NSS}	0,01	2 q_{2_NSS}
240	q_{1_NSS}	0,01	0

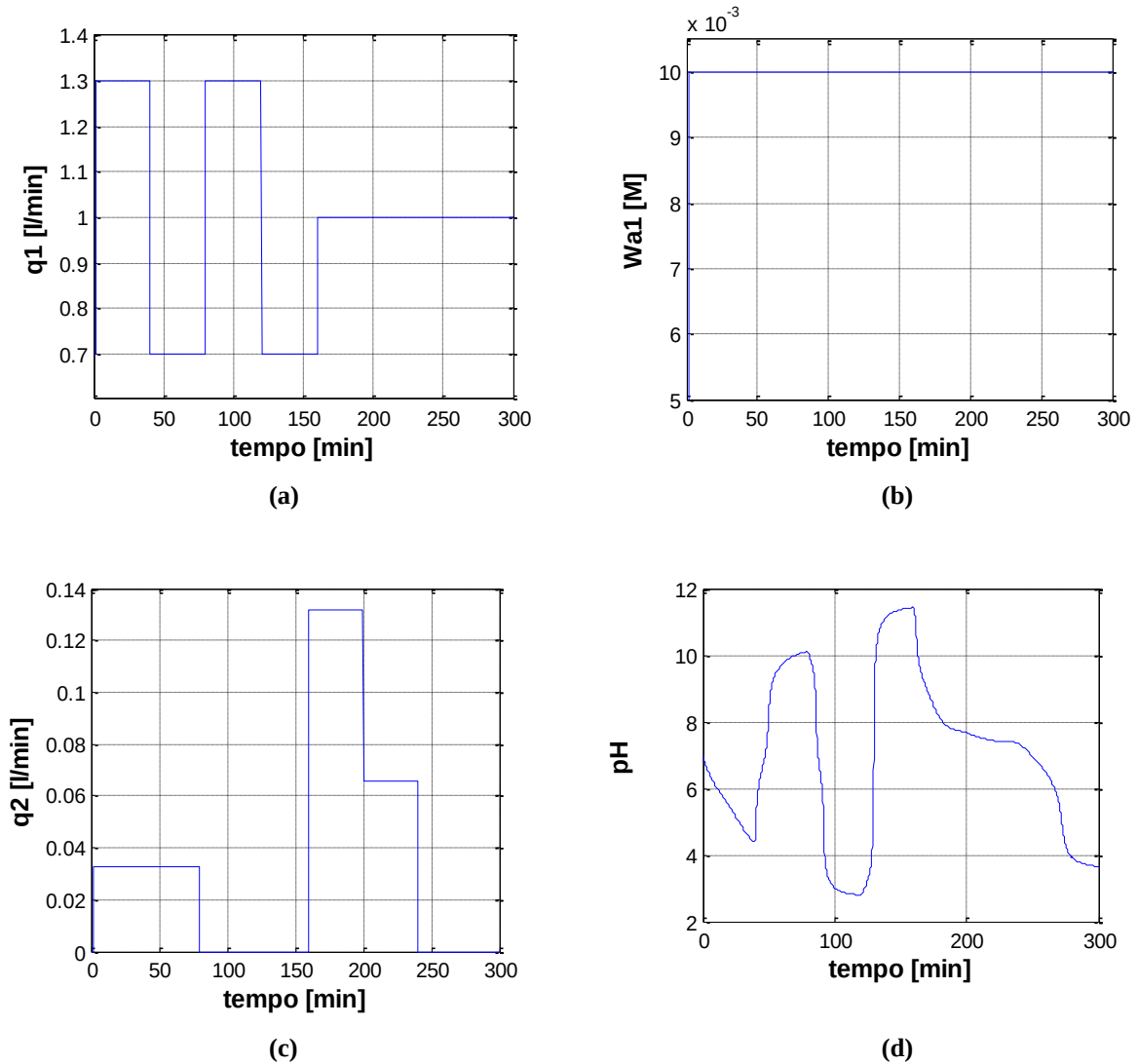
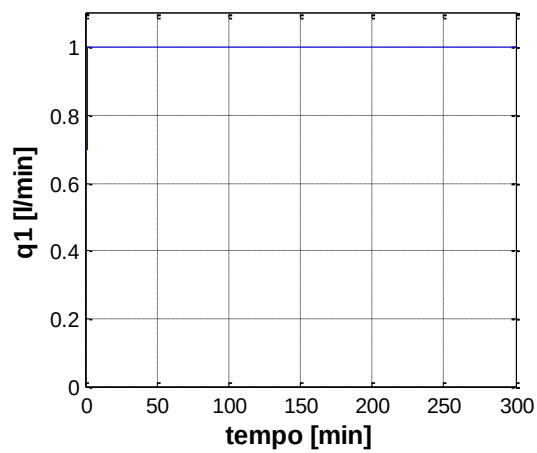


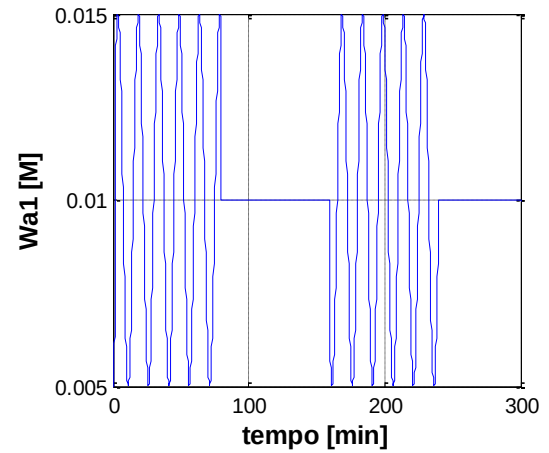
Figura 24 – Resposta de malha aberta para o teste 1: (a) vazão de ácido; (b) concentração de ácido; (c) vazão de *buffer*; (d) pH do fluxo de saída.

Tabela 5 – Teste 2: onda senoidal com período de 15 min W_{a1} com/sem *buffer*

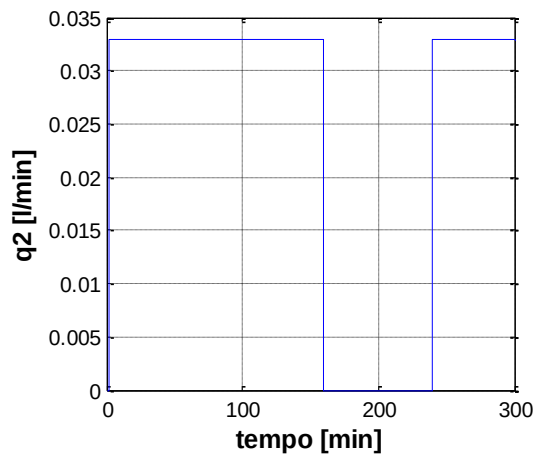
tempo (min)	q_1 (l/min)	$W_{a1}(M)$	q_2 (l/min)
0	q_{1_NSS}	$\text{Seno}(\omega t)$	q_{2_NSS}
80	q_{1_NSS}	0,01	q_{2_NSS}
160	q_{1_NSS}	$\text{Seno}(\omega t)$	0,05 q_{2_NSS}
240	q_{1_NSS}	0,01	q_{2_NSS}



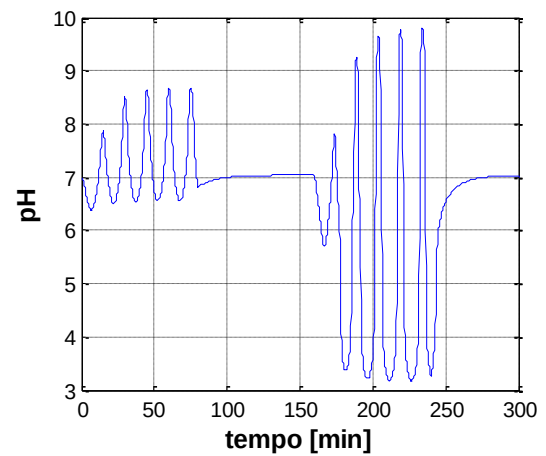
(a)



(b)



(c)



(d)

Figura 25 – Resposta de malha aberta para o teste 2: (a) vazão de ácido; (b) concentração de ácido; (c) vazão de *buffer*; (d) pH do fluxo de saída.

Como foi possível observar, as Figuras 24 e 25 representam fielmente os resultados apresentados nos artigos de Hall (HALL; SEBORG, 1989) e Alvarez (ALVAREZ et al., 2001), mostrando assim que o modelo de referência ou modelo virtual do processo de neutralização de pH, e consequentemente a sua implementação e representação usando os *softwares* Matlab®/Simulink®, poderá ser utilizado para finalidades de análise, simulação, identificação e controle deste processo, assim como gerar os dados para testar o desempenho do sensor virtual desenvolvido neste trabalho.

Assim, o uso de modelos não-lineares na análise e identificação de sistemas é, em muitos casos, essencial, devido a um número de fenômenos observados na prática, os quais resultam das não-linearidades e variância no tempo do sistema. Muitos processos industriais importantes, incluindo coluna de destilação, reação química altamente exotérmica, neutralização de pH e outros grupos de sistemas apresentam procedimento altamente não-linear e variante no tempo e podem operar em uma faixa ampla de condições de operação.

As crescentes exigências de qualidade e quantidade colocadas para a indústria de processo a defronta com situações de operações extremas, onde os efeitos não-lineares e variantes no tempo são muito mais importantes (JERÔNIMO, 2004) (SANTOS, 2007) (CASILLO, 2009).

5 RESULTADOS E DISCUSSÃO

5.1 Aplicações do método proposto

Com o intuito de testar o desempenho da metodologia proposta no Capítulo 3, ou seja, a união do método de agrupamentos nebulosos de Gustaffson-Kessel (BABUŠKA; VERBRUGGEN, 1997), para tratar a parte antecedente das regras do algoritmo nebuloso, em conjunto com o método proposto por Wang-Langari (WANG; LANGARI, 1996b), que propuseram uma forma mais simples de variação do fator de esquecimento para ser utilizado na identificação dos parâmetros da parte consequente das regras nebulosas, foi utilizado um exemplo. Esse exemplo, uma aplicação de algoritmos de identificação recursiva de sistemas, quando os parâmetros do sistema são variantes no tempo e onde é absolutamente necessário acompanhar a variação dos parâmetros em tempo real, é o sistema inicialmente proposto por Moustafa (1983) e também utilizado por Wang e Langari (1994) (1996b) e descrito a seguir.

Considerando um sistema com uma entrada e uma saída descrito pela equação a diferença de primeira ordem:

$$y_k = a_k y_{k-1} + b_k u_{k-1} + \omega_k \quad (63)$$

onde a_k e b_k são dois parâmetros variantes no tempo com valores iniciais $a_1 = 0,8$, e $b_1 = 2,0$; e o sinal de entrada u_k é uma função senoidal $u_k = 0,7\sin(k)$ e ω_k é um ruído Gaussiano com média zero e variância 0,7. As seguintes constantes foram escolhidas para o fator de esquecimento variável:

$$\lambda_{\min} = 0,2, \quad \lambda_{\max} = 1, \quad \Delta = 0,995 \quad \Gamma = 1000$$

Foram considerados dois tipos de variação dos parâmetros, uma suave e outra abrupta.

No caso de mudança suave nos valores dos parâmetros, a_k e b_k sofrem variação de acordo com as dinâmicas

$$a_k = \begin{cases} \left(1 - \frac{1}{k^2}\right) a_{k-1}, & 2 \leq k \leq 500 \\ \left(1 - \frac{1}{k}\right) a_{k-1}, & k \geq 501 \end{cases}$$

$$b_k = \begin{cases} \left(1 - \frac{1}{k^2}\right) b_{k-1}, & 2 \leq k \leq 500 \\ \left(1 - \frac{1}{k}\right) b_{k-1}, & k \geq 501 \end{cases}$$

Foram geradas 1000 amostras de dados simulados da Equação (63) e construído um modelo nebuloso, usando para isso a metodologia descrita no Capítulo 3, ou seja, para tratamento das variáveis de entrada, parte antecedente das regras, utiliza-se agrupamento nebuloso de Gustafson-Kessel e para estimação dos parâmetros da parte conseqüente utiliza-se o conhecido mínimos quadrados recursivo, mas com fator de esquecimento variável.

Com esse resultado é possível mostrar a habilidade do sensor virtual para identificação paramétrica, o mesmo sendo usado como um estimador de parâmetros de sistemas variantes no tempo. A Figura 26 apresenta os dados de entrada e saída.

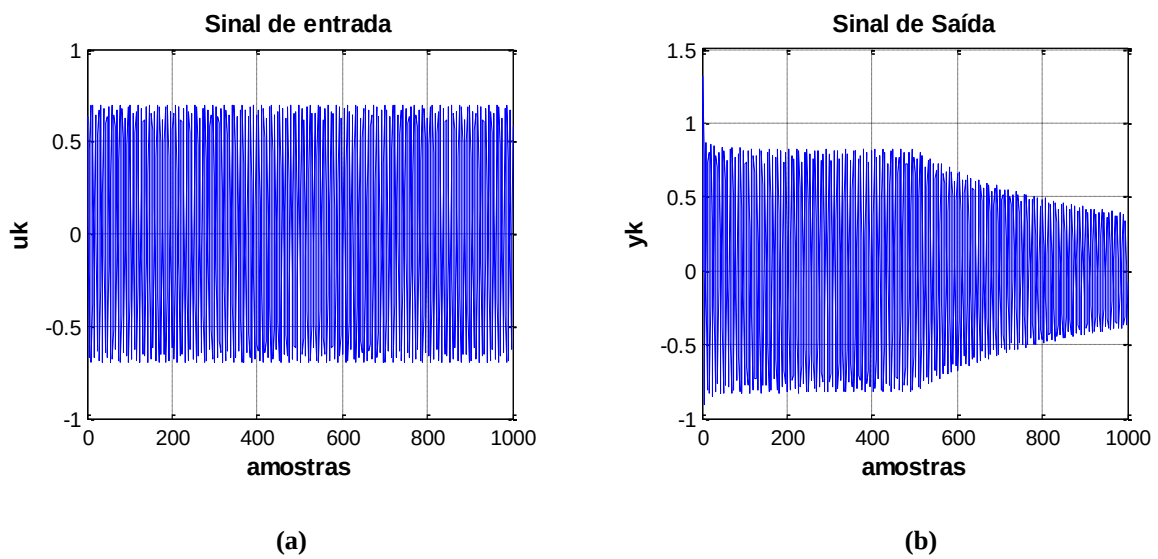
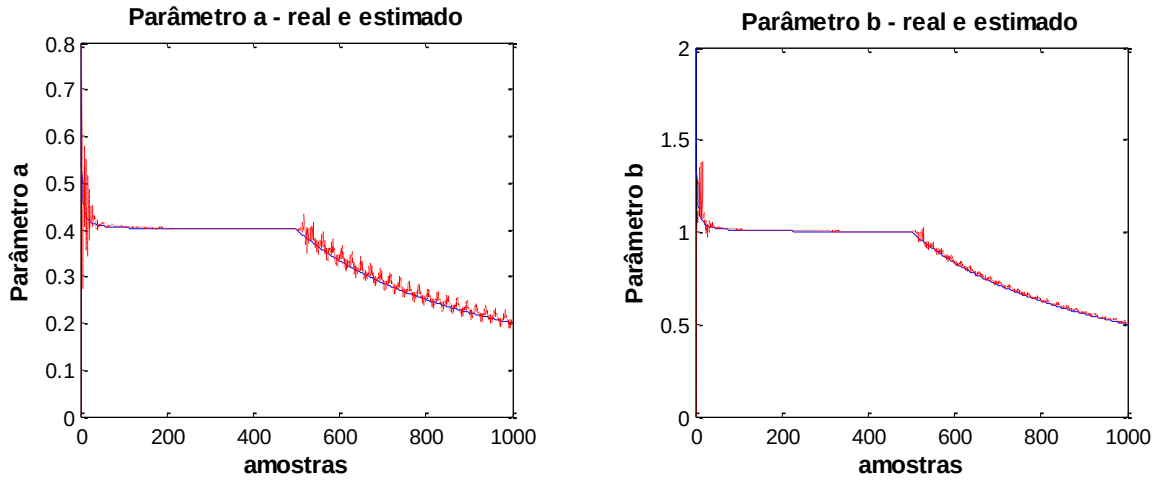
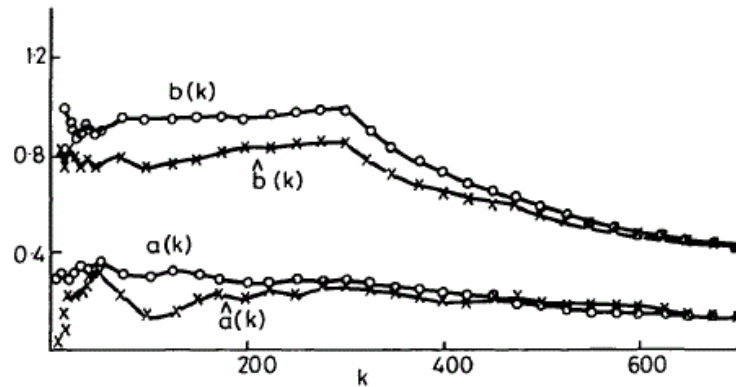


Figura 26: Sinais de entrada (a) e saída (b) da Equação (63) com variação lenta dos parâmetros.

A Figura 27(a) mostra os parâmetros a e b estimados pelo modelo nebuloso proposto neste trabalho. É possível observar que os parâmetros rapidamente convergem para o valor esperado e quando ocorre a variação paramétrica, apesar de ser lenta, consegue seguir a tendência desta variação mantendo seu valor médio próximo do valor real



(a)

Real-time tracking of $a(k)$ and $b(k)$ with sinusoidal input

(b)

Figura 27: (a) Parâmetros a e b , real (---) e estimado (---). (b) Imagem retirada de Moustafa (1983)

A título de comparação do método utilizado com outro método para identificação de sistemas variantes no tempo, é apresentado na Figura 27(b), retirada do trabalho de Moustafa (1983), os resultados por ele obtidos, para os parâmetros a e b . Moustafa foi

o primeiro a utilizar este exemplo de aplicação, porém o método de identificação utilizado em seu trabalho é o pouco conhecido funções de Lyapunov estocásticas. Neste mesmo trabalho, o autor relata em sua conclusão a “quase certa convergência do erro de identificação para um valor nulo, para o caso onde o ruído decorrente da dinâmica dos parâmetros possui uma variância adicional”.

Observa-se que a identificação nebulosa proposta responde bem com relação à convergência das estimativas paramétricas e, mesmo quando ocorre uma variação lenta porém constante nos parâmetros, o erro é muito pequeno, ao passo que no trabalho de Moustafa (1983) ele apenas comenta em sua conclusão, mas não mostra em simulação ou mesmo analiticamente, este resultado. Em suma, o algoritmo desenvolvido nesta tese é superior àquele de Moustafa (1983), uma vez que tem convergência mais rápida e menor erro de estimação.

A Figura 28(a) apresenta a variação do fator de esquecimento λ e a Figura 28(b) apresenta o valor do erro de saída entre o valor simulado e o estimado pelo modelo.

Como se pode observar na Figura 28(a), o método dos mínimos quadrados recursivos com fator de esquecimento variável, usando a variação proposta por Wang e Langari (1996b), adapta os valores do fator de esquecimento, no cálculo dos parâmetros na parte consequentes das regras, com o propósito de se adaptar à variação ocorrida nos mesmos de forma a minimizar o erro (Figura 28(b)).

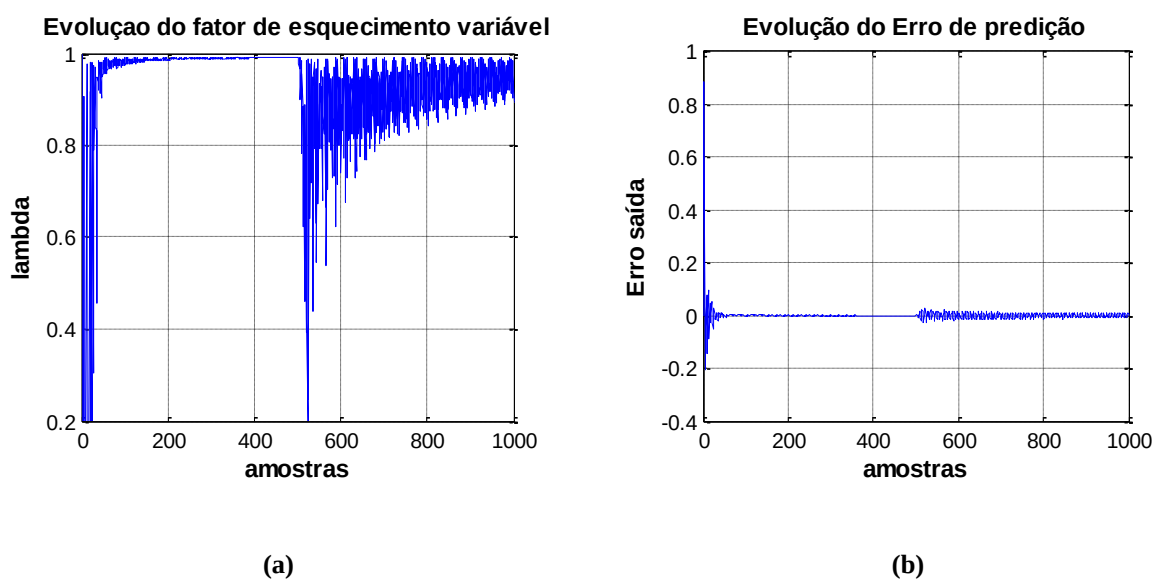


Figura 28: (a) Fator de esquecimento variável, (b) Erro de saída entre o valor simulado e o estimado.

No caso de mudança abrupta nos valores dos parâmetros, a_k e b_k variam de acordo com as dinâmicas

$$a_k = \begin{cases} 0,8, & 2 \leq k \leq 500 \\ 0,4, & k \geq 501 \end{cases}$$

$$b_k = \begin{cases} 2,0, & 2 \leq k \leq 500 \\ 1,0, & k \geq 501 \end{cases}$$

Também foram geradas 1000 amostras de dados da Equação (63), e então construído o modelo nebuloso proposto neste trabalho. O Anexo B contém um *script-file* em Matlab para solução do exemplo em questão. A Figura 29 apresenta os dados de entrada e saída.

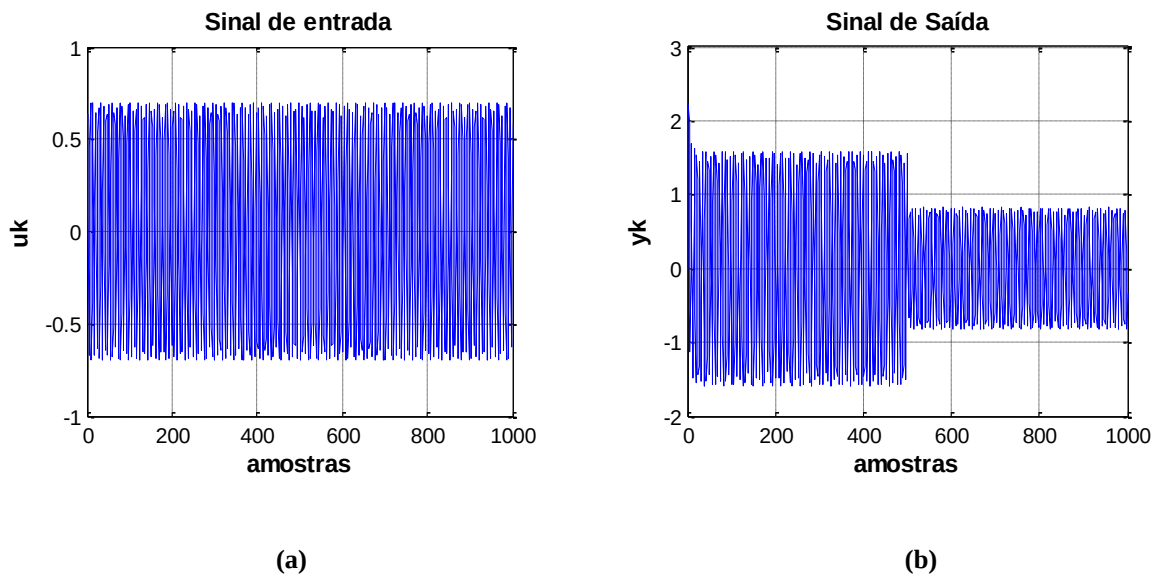


Figura 29: Sinais de entrada (a) e saída (b) da Equação (63) com variação abrupta dos parâmetros.

Agora o modelo nebuloso estimado, usando a metodologia proposta nesta tese, apresenta como parâmetros os valores mostrados na Figura 30. Como se observa nas Figuras 30(a) e 30(b), o algoritmo proposto é capaz de rapidamente identificar os parâmetros do modelo, mesmo em condições de mudança repentina nestes.

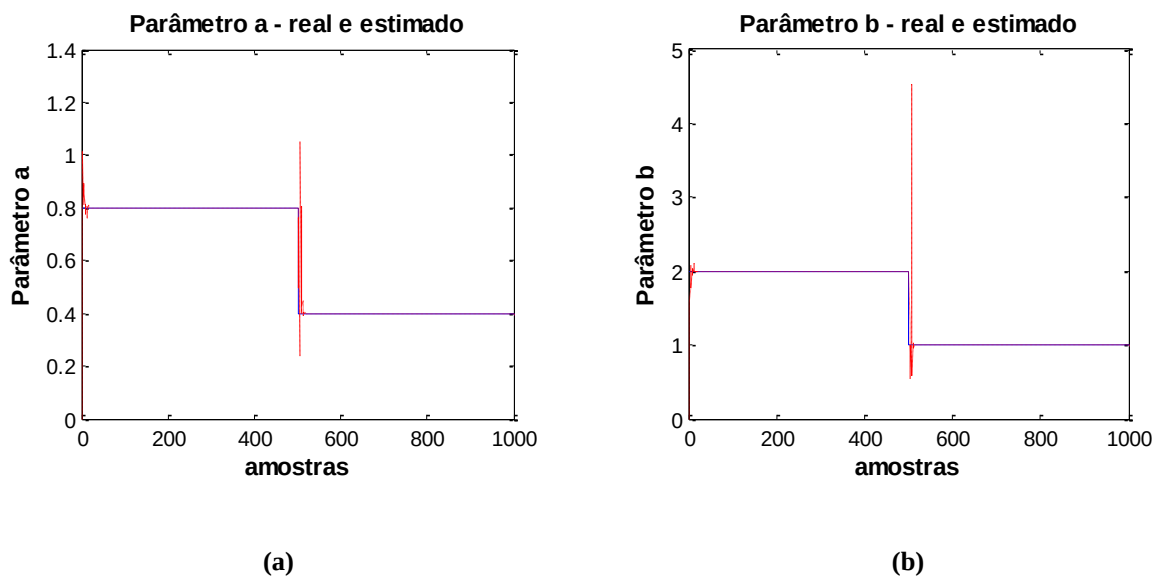


Figura 30: Parâmetros a e b , real (---) e estimado (---).

A Figura 31(a) apresenta a variação do fator de esquecimento λ , e a Figura 32(b) apresenta o valor do erro de saída entre o valor simulado e o estimado pelo modelo. Mais uma vez, a Figura 31(a) ilustra o esforço do modelo proposto neste trabalho, alterando o valor do fator de esquecimento, com o intuito de minimizar o erro de predição, garantindo também que as estimativas não polarizem, ou seja, tendam para um valor diferente do real. Observa-se que rapidamente o erro aproxima-se de zero.

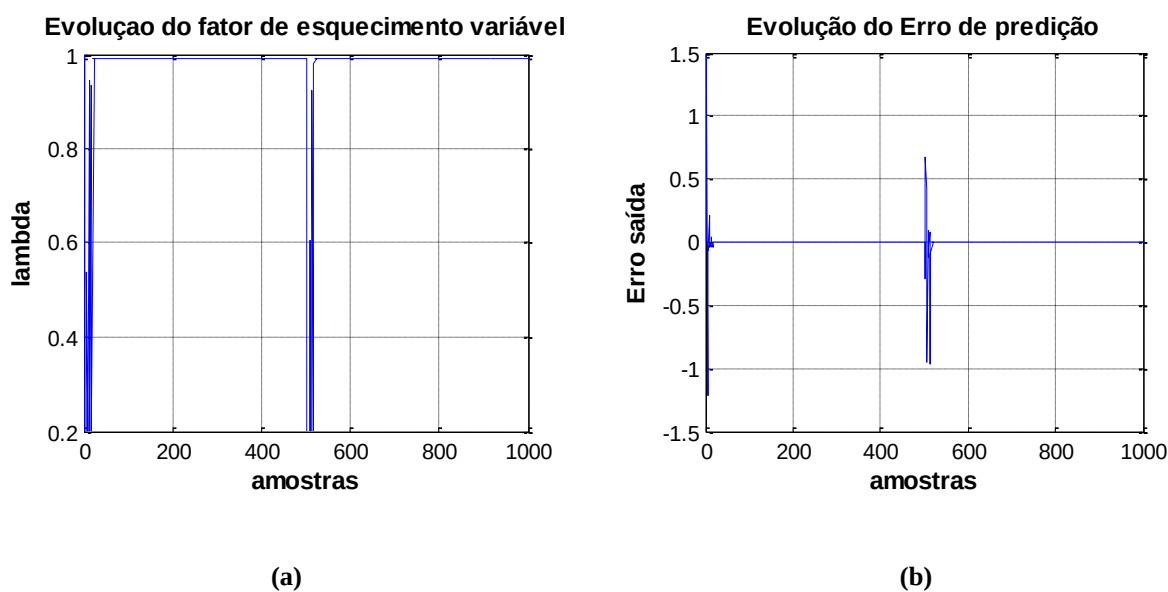


Figura 31: (a) Fator de esquecimento variável, (b) Erro de saída entre o valor simulado e o estimado.

Tanto para o sistema com variação lenta como para aquele com variação abrupta dos parâmetros, foram escolhidos y_{k-1} e u_{k-1} como variáveis de entrada do modelo nebuloso, e, usando a técnica de discretização nebulosa (WANG; LANGARI, 1996a), foi determinado 6 (seis) como número de regras utilizado. Vale ressaltar que, neste mesmo artigo, Wang e Langari apontam como um problema no uso da técnica de discretização nebulosa proposta por eles, técnica esta baseada em uma modificação do método de clusterização *fuzzy c-means* para identificar a estrutura dos antecedentes, o fato de que ao invés de calcular a função de pertinência de cada variável de entrada *on-line* estas funções são calculadas *off-line*, apresentando assim uma certa dificuldade para tratar sistemas variantes no tempo. Com isso, a abordagem proposta por eles para tratar a parte consequente das regras nebulosas trabalha muito bem com sistemas variantes no tempo, mas a parte antecedente apresenta dificuldades. Portanto, a utilização de um método que trate a parte antecedente das regras da forma como é proposta por Gustafson e Kessel (GUSTAFSON; KESSEL, 1978) pode resolver o problema da clusterização dos antecedentes de forma satisfatória.

Pode ser observado nas Figuras 27 e 30, que mostram, respectivamente, os parâmetros reais e os estimados do sistema com variação lenta e variação abrupta, que os dois algoritmos em conjunto, o de agrupamento nebuloso pelo método de Gustafson-Kessel (GUSTAFSON; KESSEL, 1978) (BABUŠKA; VERBRUGGEN, 1997) e o do fator de esquecimento variável de Wang-Langari (WANG; LANGARI, 1996a), podem seguir os parâmetros variantes no tempo muito bem, ou seja, uma excelente habilidade de rastreamento, o que, em conjunto, são uma ótima opção para tratar sistemas variantes no tempo. A extensa revisão bibliográfica feita nesta tese mostra que esta abordagem conjunta está sendo inaugurada neste trabalho.

Um outro exemplo de aplicação da metodologia para sensores virtuais proposta nesta tese foi retirado do trabalho de Narendra e Parthasarathy (1990), mostrando um dos resultados simulados, enquanto os demais resultados são utilizados para teste do controlador estudado por eles, em que a planta a ser identificada é dada pela seguinte equação a diferença de segunda ordem e não-linear

$$y_k = \frac{y_{k-1}y_{k-2}(y_{k-1} + 2,5)}{1 + y_{k-1}^2 + y_{k-2}^2} u_{k-1}. \quad (64)$$

700 amostras de dados simulados são gerados a partir da equação 64; as primeiras 500 amostras são obtidas utilizando como entrada um sinal aleatório uniformemente distribuído no intervalo de $[-2, 2]$, e as 200 amostras restantes são obtidas através de um sinal senoidal na forma $u_k = \text{sen}(2\pi k / 25)$.

Foram utilizadas as primeiras 500 amostras para construir o modelo nebuloso usando a metodologia proposta nesta tese. O desempenho do modelo nebuloso é testado usando as 200 amostras restantes. Foram selecionadas como variáveis de entrada y_{k-1} , y_{k-2} e u_{k-1} , o número de agrupamentos ajustado para 3, devido a uma limitação de *hardware* para simulação. Todas as simulações foram feitas em um microcomputador PC de 800MHz com 512MB de memória RAM.

A Figura 32(a) mostra o sinal de entrada aleatório (500 amostras iniciais), e a Figura 32(b) mostra o sinal de entrada usado na validação do modelo, que pode ser considerado como sinal de teste em uma condição de predição do modelo.

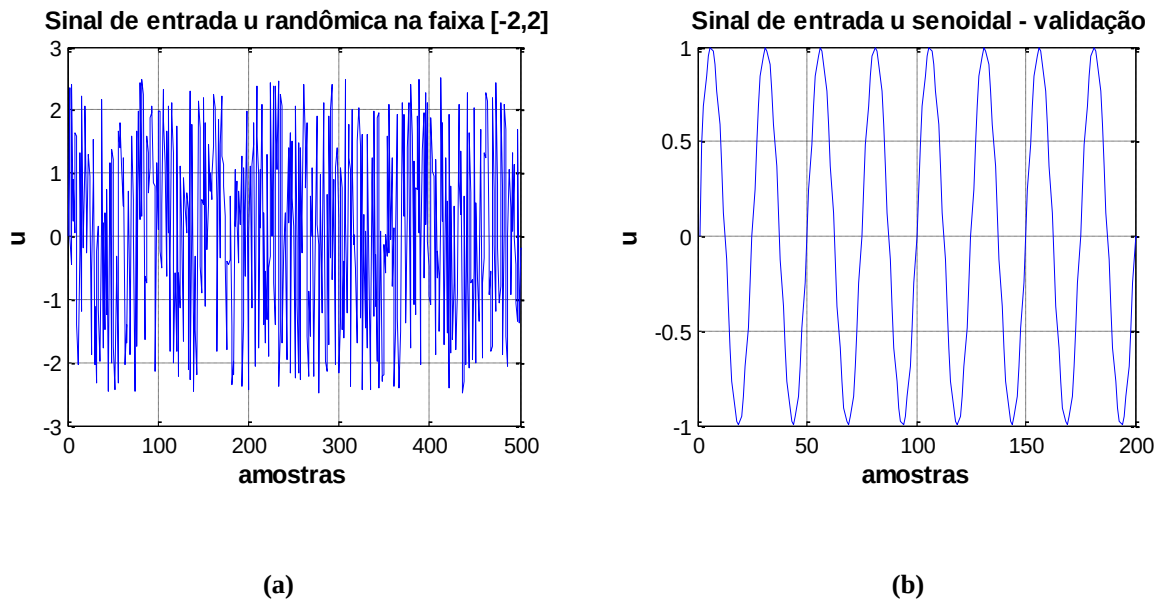


Figura 32: (a) Sinal de entrada aleatório, (b) sinal de entrada periódico.

É mostrado na Figura 33(a) o sinal de saída simulado em conjunto com o sinal de saída do modelo estimado, usado na fase de treinamento, e na Figura 33(b), o sinal de saída simulado e o estimado, usado para validação do modelo.

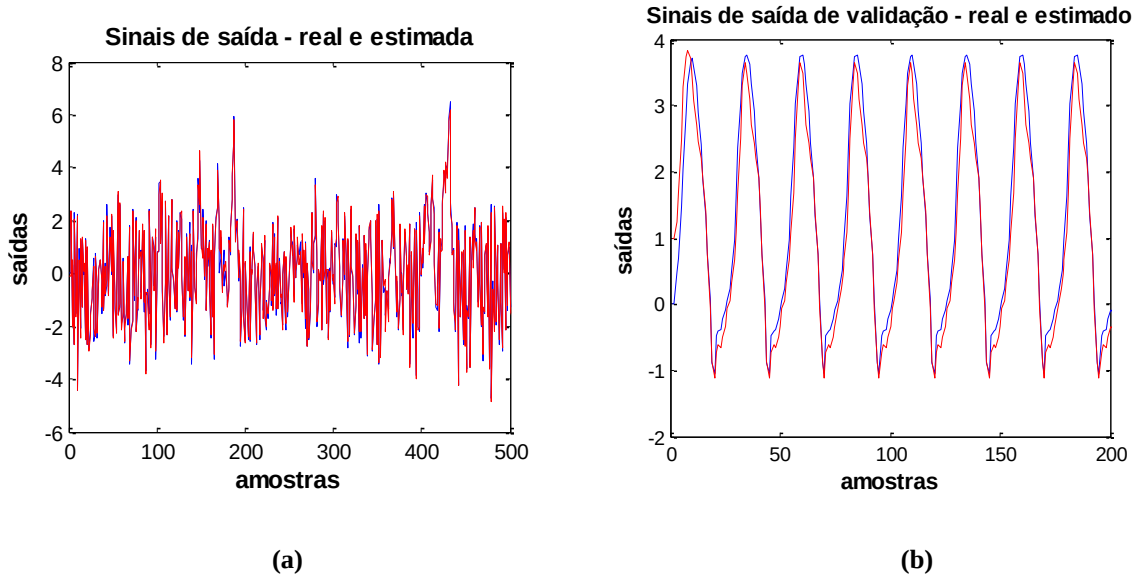


Figura 33: (a) Sinais de saída real (---) estimado (---) - treinamento, (b) Sinais de saída real (---) estimado (---) - validação.

É possível notar que o modelo apresenta um pequeno erro médio quadrático, e_k , calculado pela equação (65) e apresentado na Figura 34, juntamente com o erro de predição. y_k = valor simulado ; $\hat{y}_k$ = valor estimado

$$\text{RMS}(e_k) = \left(\frac{1}{k} \sum_{i=1}^k e_i^2 \right)^{1/2}, \quad e_k = y_k - \hat{y}_k \quad (65)$$

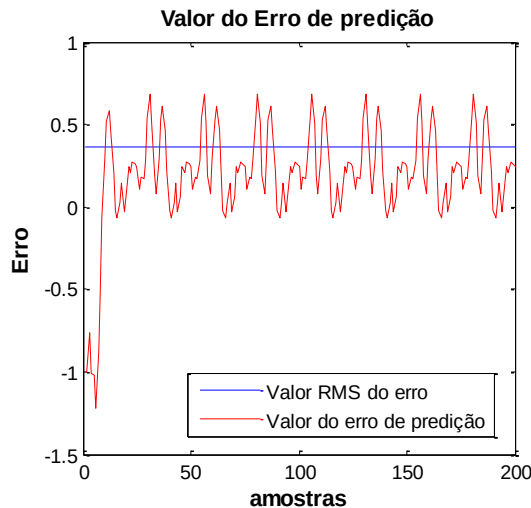


Figura 34: Valor RMS do erro e valor do erro de predição.

O erro de estimação mostrado poderia ser diminuído, se fosse utilizado um número maior de agrupamentos, pois por limitações de hardware utilizaram-se apenas 3 agrupamentos, como descrito anteriormente.

Wang e Langari (1996) utilizaram o mesmo exemplo para testar seu algoritmo, mas não mostraram uma comparação entre o dado simulado e o dado estimado. Além disso, eles utilizaram 8 agrupamentos e, mesmo assim, lembram em sua conclusão que utilizaram 500 amostras para construir o modelo nebuloso, enquanto Narendra e Parthasarathy (1990) utilizaram nada menos que 100.000 amostras para identificar um modelo usando redes neurais. Finalizam afirmando que “é esperado que o desempenho do modelo nebuloso identificado possa ser melhorado se o número de amostras usadas para construção do modelo for aumentado”.

Foi feita a verificação do desempenho do modelo nebuloso identificado com o aumento do número de amostras. A Tabela 6 apresenta o valor RMS do erro nos dados simulados. Com as mesmas condições de entradas anteriores, não foi possível observar grandes melhorias.

Tabela 6 – Variação do erro RMS em função do número de amostras de treinamento

n° amostras	500	1000	2000	3000	4000	5000
e_{RMS}	0,4588	0,3515	0,3437	0,3755	0,3924	0,3613

Contudo, com relação ao sinal de saída simulado e o sinal de saída estimado, nos dados validados, é possível observa-se uma sensível diferença, como é mostrado na Figura 35, quando se usa um número maior de amostras na fase de treinamento do estimador (sensor virtual, no contexto desta tese). Neste exemplo, usaram-se 2.000 amostras:

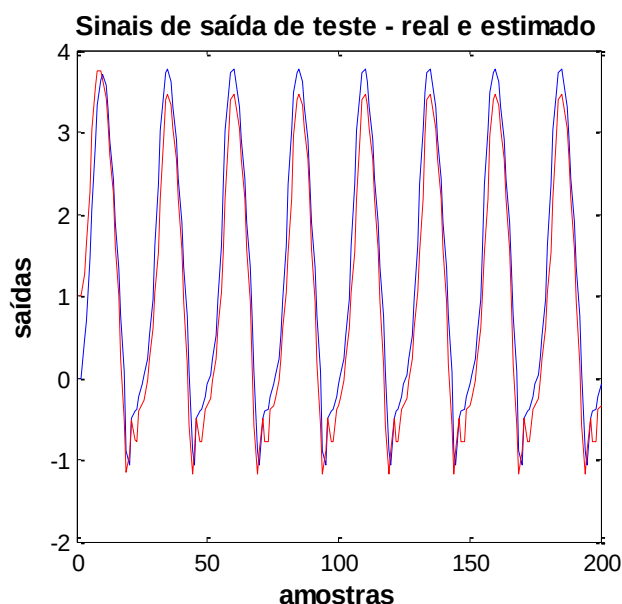


Figura 35: Sinais de saída real (---) estimado (---) - validação.

5.2 Sensor virtual aplicado ao processo de neutralização de pH

O exemplo de simulação a seguir ilustra a aplicação do sensor virtual, utilizando o método proposto, ao processo de neutralização de pH. Este processo, altamente não linear, é muito conhecido por colocar em posição de dificuldade muitas técnicas de identificação caixa-preta. A Figura 36 mostra um esquema utilizado no modelo virtual deste processo, descrito no Capítulo 4, e mostrado novamente aqui de forma resumida. Esta abordagem poderia ser considerada do tipo caixa-preta, apesar da descrição mais aprofundada do processo químico ter sido realizada anteriormente.

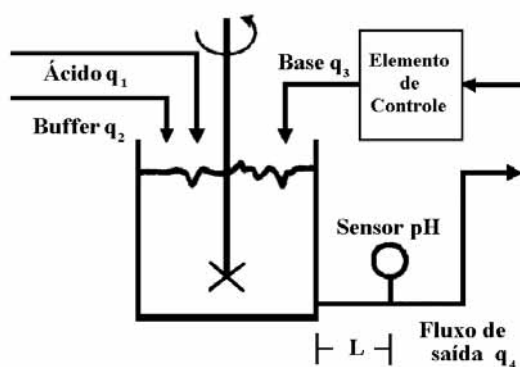
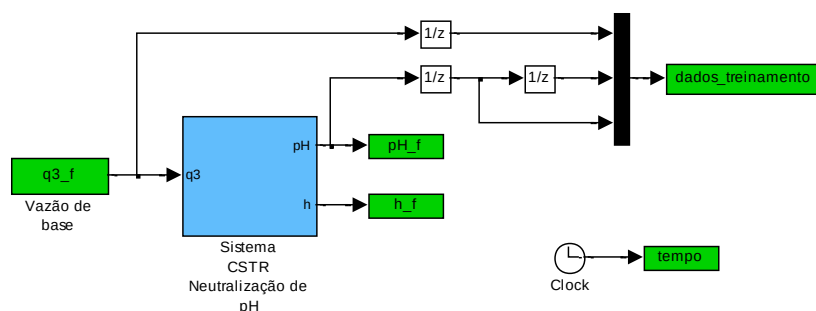


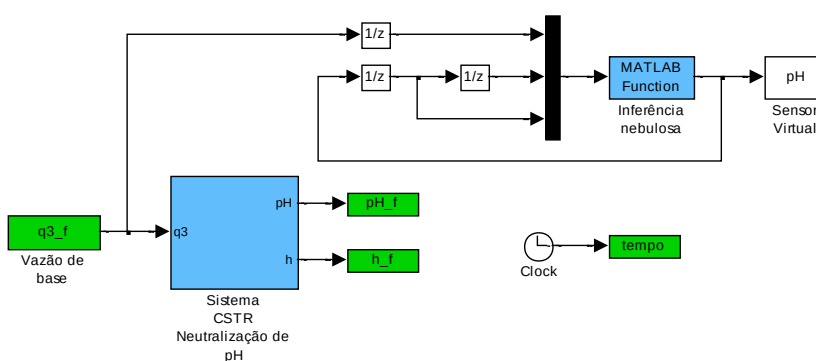
Figura 36: Modelo simplificado do tanque de neutralização de pH utilizado.

Na simulação, atrasadores são utilizados na obtenção dos dados necessários ao treinamento. Estes dados são agrupados e armazenados na memória como uma matriz.

O esquema de simulação em Simulink®, realizado no modelo virtual desenvolvido no Capítulo 4 e mostrado na Figura 37(a), é utilizado na fase de treinamento e mostra uma configuração de coleta de dados para 1 regressor de vazão e dois de pH. Procurou-se aqui utilizar os mesmos dados propostos no artigo de Babuška e Verbruggen (1997), com uma diferença: neste artigo, eles propõem identificar este processo usando a técnica de agrupamentos nebulosos de Gustafson-Kessel, porém nos consequentes das regras para identificação nebulosa eles utilizam um outro método de estimação paramétrica, que eles chamaram de “mínimos quadrados totais”, feito de forma *off-line*.



(a)

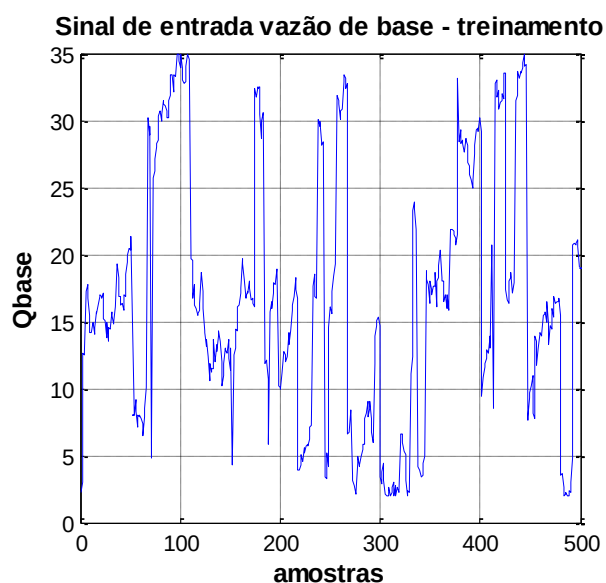


(b)

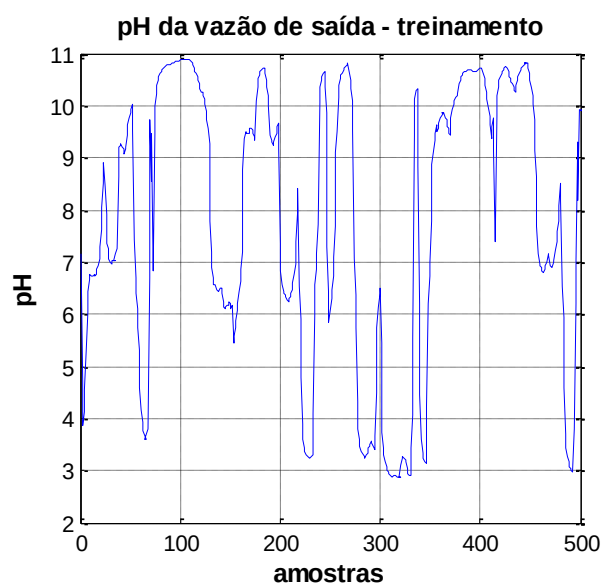
Figura 37: (a) esquema de simulação usado no treinamento, (b) esquema usado como sensor virtual.

Após o treinamento, a validação é executada com a configuração da Figura 37(b). Observa-se que, após o treinamento, as medidas de pH não são mais utilizadas na

estimativa de um novo valor, como é desejado. A Figura 38(a) mostra o sinal de entrada do processo CSTR, ou seja, a vazão de base q_3 . O conjunto de dados para o treinamento contém 500 amostras, com um intervalo de amostragem de 15s. A Figura 38(b) mostra o sinal de saída, o valor de pH. Com base em conhecimentos anteriores, como descrito no Capítulo 4 desta tese, esta estrutura é considerada apropriada para aproximações da dinâmica do pH.



(a)



(b)

Figura 38: (a) entrada vazão de base q_3 , (b) saída pH medido.

A Figura 39 apresenta uma comparação entre o pH medido, ou seja, o pH simulado com os dados estimados na fase de treinamento, e o pH estimado, sendo possível observar uma boa capacidade de estimação do pH.

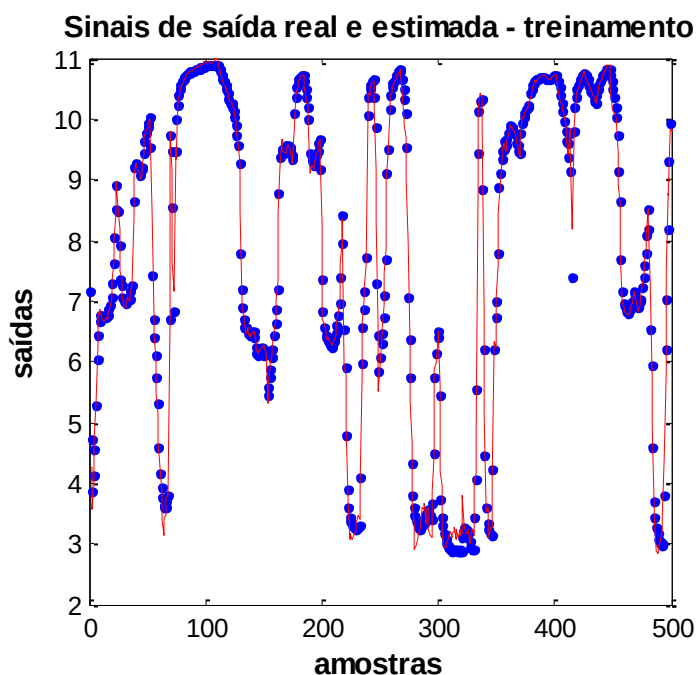


Figura 39: Sinais de saída real (•••) e estimado (---) - treinamento.

Na fase de validação, que corresponde à fase onde o algoritmo trabalha como sensor virtual, são utilizadas mais 500 amostras para os sinais de entrada de vazão, e as medidas de pH não são mais utilizadas na estimativa de um novo valor, como era desejado.

A Figura 40(a) mostra os sinais de saída real e o sinal de saída do sensor virtual. É possível observar um resultado bastante coerente, e que pode ser confirmado comparando este resultado com o da Figura 40(b), que é uma figura transcrita do artigo de Babuška e Verbruggen (1997) mostrando o resultado da validação do modelo nebuloso proposto por eles. Chama-se a atenção para o fato de que a identificação proposta por eles não serviria para ser utilizada em tempo real, ou seja, como um sensor virtual, pois ela foi desenvolvida como um método de identificação nebulosa de sistemas complexos, operando apenas *off-line*, e sem a capacidade de adaptação em tempo real (usa-se o LS total no cálculo do conseqüente).

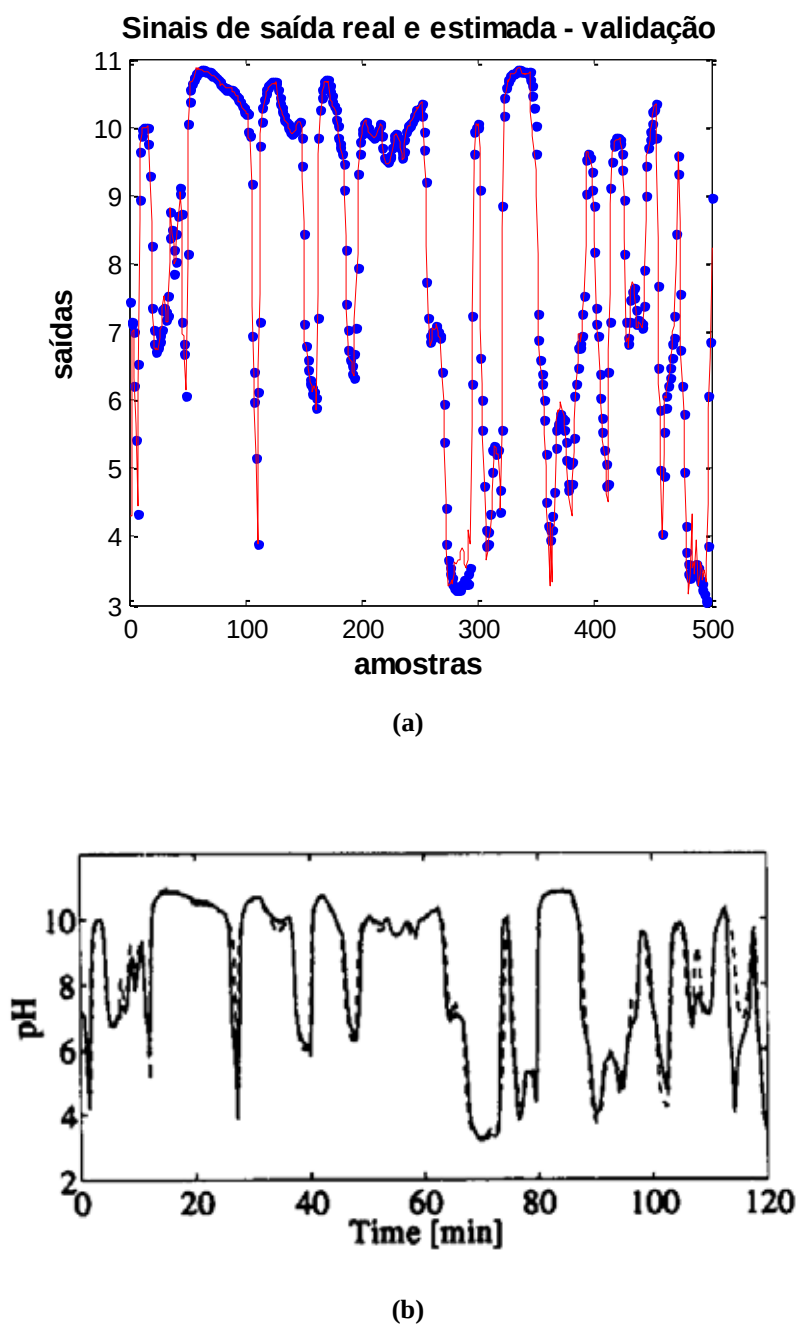


Figura 40: (a) Sinais de saída real (•••) e estimado (---) – validação, (b) resultado obtido por Babuška e Verbruggen (1997).

A Figura 41 apresenta o erro de predição do sensor virtual de pH proposto neste trabalho. É possível notar que, apesar de ser diferente de zero, o erro possui valor médio muito próximo de zero, permitindo o sinal estimado rastrear o sinal real, com desempenho aceitável para uma gama de aplicações industriais/científicas.

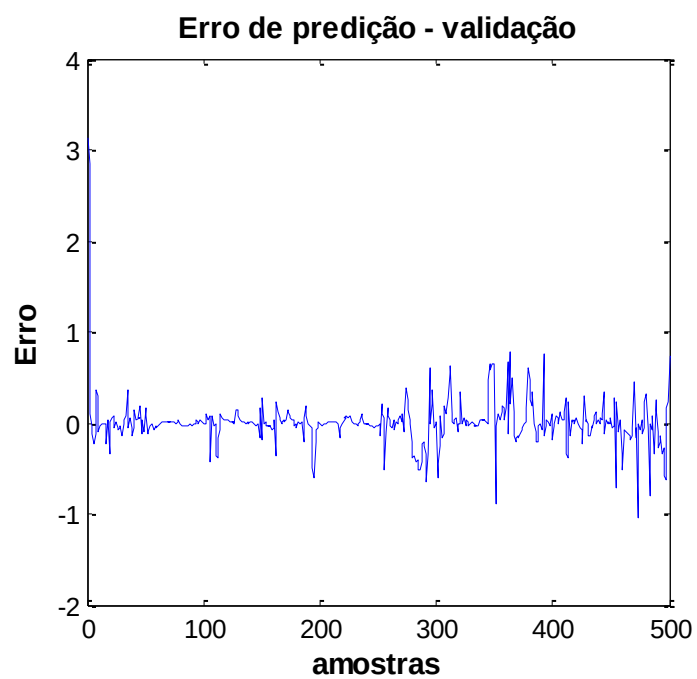


Figura 41: Valor do erro de predição.

6 CONCLUSÃO E CONSIDERAÇÕES FINAIS

O desenvolvimento de um sensor virtual para variáveis que não podem ser medidas diretamente por um sensor físico representa um avanço no controle de processos. Dificuldades encontradas no emprego da tecnologia atual, como a inexistência de sensores, preços elevados, atrasos ou imprecisão das medidas, podem, com o emprego de sensores virtuais, serem contornadas.

Durante a pesquisa bibliográfica, foi possível constatar que os métodos de modelagem ou identificação de sistemas complexos, ou seja, principalmente sistemas não-lineares, eram basicamente divididos em dois grandes grupos:

- Os métodos que descrevem o sistema usando relações funcionais não-lineares entre as variáveis deste sistema, conhecidos como métodos Globais; alguns exemplos são: os modelos de espaço de estado não-linear, modelos caixa-preta de entrada e saída NARX em conexão com redes neurais e wavelets.

Muitos exemplos de aplicações reais foram encontrados no livro de Fortuna et al. 2007, e nos vários artigos de Gonzalez, entre os anos de 1994 a 2005. Um fato curioso, apesar das importantes aplicações descritas por Gonzalez, em periódicos e congressos do IEEE, é que não se encontra nenhuma citação sobre seus trabalhos e contribuições no livro de Fortuna a respeito da utilização de métodos Globais em identificação de sistemas complexos e desenvolvimento de sensores virtuais.

- Métodos que tentam lidar com a complexidade e não-linearidade de sistemas, através da decomposição do problema de modelagem em um número maior de problemas mais simples, em muitos casos, sub-problemas lineares, conhecidos como métodos Locais. Estes métodos são conceitualmente simples e atraentes e, com isso, são geralmente mais facilmente interpretáveis que os complicados modelos globais. Técnicas de modelagem ou identificação baseadas em conjuntos nebulosos são exemplos de métodos de modelagem Local.

Em 1973, Zadeh apresentou, como parte de uma importante formulação sua, o princípio da incompatibilidade:

“À medida que a complexidade de um sistema aumenta, nossa habilidade para fazer afirmações precisas e que sejam significativas acerca deste sistema diminui até que um limiar é atingido além do qual precisão e relevância tornam-se quase que características mutuamente exclusivas”.

O objetivo desta pesquisa é o desenvolvimento de sensores virtuais para monitoramento de variáveis do processo, ou seja, processos complexos não-lineares e variantes no tempo, o que ainda não havia sido realizado até então, como mostra a literatura técnico-científica.

As técnicas de modelagem e identificação baseadas em conjuntos nebulosos surgiram como técnicas promissoras, pois, através de alguns trabalhos de Wang e Langari (1996), observou-se a utilização das mesmas em identificação de sistemas complexos não-lineares e variantes no tempo; porém, elas apresentavam algumas limitações na identificação em tempo real, o que é importante no desenvolvimento de sensores virtuais.

Nesta mesma época, e de forma independente, Babuška e Verbruggen (1997) propuseram um método de identificação de sistemas complexos não-lineares e variantes no tempo que eles empregaram com sucesso, porém novamente não utilizado em tempo real.

O mais interessante é que a parte limitante da técnica proposta por Wang e Langari, ou seja, a que não permitia ou limitava seu uso em tempo real, era a que demonstrava ser a grande solução ou contribuição apresentada por Babuška e Verbruggen e vice-versa, ou seja, a limitação encontrada no algoritmo de Babuška e Verbruggen foi resolvida muito bem por Wang e Langari.

Mas por que isso ocorreu? A resposta, apesar de ser simples e possuir uma implementação deste algoritmo não trivial, é que tanto Wang e Langari quanto Babuska e Verbruggen estavam trabalhando, ou melhor, procurando desenvolver técnicas de identificação e modelagem baseadas em conjuntos nebulosos para sistemas complexos com a finalidade de obter um modelo para o processo e não a criação ou desenvolvimento de um sensor virtual.

A partir de informações coletadas de um processo específico, é possível, através dos algoritmos discutidos, chegar-se a um modelo nebuloso relativamente preciso em

uma abordagem do tipo caixa-preta. Além disso, o método de agrupamento de Gustafson-Kessel (*product space clustering*) mostrou-se versátil o suficiente para permitir a geração de modelos de processos muito mais complexos e com um número muito maior de variáveis de entrada do que o permitido pelo método original de Wang e Langari. Assim, o uso de técnicas de identificação de modelos nebulosos mostrou-se relativamente eficiente na geração de sensores virtuais.

Uma das dificuldades da modelagem de um sistema em tempo discreto utilizando a teoria nebulosa é a determinação da estrutura do modelo nebuloso, ou seja, a determinação do número de regressores de cada variável, necessários à geração dos modelos. O uso de um número pequeno de regressores das variáveis de entrada muitas vezes gera um modelo impreciso, enquanto o uso de excessivos regressores exige um esforço computacional muito maior. Em alguns casos, ainda é possível que alguns regressores das entradas sejam substituídos por regressores da variável estimada.

Verificou-se que a validade dos modelos nebulosos gerados é muito dependente dos dados coletados do processo real. A falta de informação destes dados gera um modelo nebuloso limitado quanto à exatidão e que, possivelmente, irá divergir, se usado na configuração de sensor virtual. Desta forma, se possível, é importante coletar os dados para diferentes excitações do sistema, de forma a prover informações sobre os diferentes estados deste.

Atrelado ao número de variáveis de entrada e aos dados de treinamento está o melhor número de agrupamentos, a ser definido antes da criação do modelo nebuloso. Uma quantidade insuficiente de agrupamentos limita o modelo quanta à capacidade de identificar as não-linearidades do processo, enquanto um número excessivamente grande de agrupamentos possivelmente modelará os erros das medidas, prejudicando a exatidão das estimativas.

Também foi observado que a capacidade de um computador gerar ou não um modelo mais complexo está relacionada ao número de regras do modelo nebuloso. Determinando-se este número máximo de regras, evita-se o início de um treinamento que horas mais tarde possa incorrer em um erro por falta de memória.

Procurou-se mostrar também neste trabalho, antes de empregar o método proposto como um sensor virtual, o desempenho do mesmo na identificação de

sistemas complexos “em tempo real”, primeiramente para sistemas variantes no tempo e, depois, para sistemas com forte não-linearidade, através de exemplos encontrados na literatura técnica, cujos resultados são muito satisfatórios nas duas condições citadas anteriormente.

Por fim, o melhor método, dentre os estudados, para a aplicação em planta de neutralização de pH, foi o Gustafson-Kessel associado aos mínimos quadrados recursivos com fator de esquecimento variável.

No entanto, é possível que, para diferentes aplicações em diferentes processos, o método de Wang e Langari, com um método de agrupamento baseado em *fuzzy c-means*, ou mesmo o agrupamento nebuloso de Gustafson-Kessel associado ao método dos mínimos quadrados, possa funcionar de forma tão satisfatória quanto o proposto nesta tese, para uso como algoritmo de identificação de sistemas complexos.

Um desdobramento deste estudo é a busca por formas para se estimar o número ótimo de agrupamentos, sem que haja necessidade de geração e validação de uma grande quantidade de modelos, ou seja, de forma quase empírica. Estas medidas, chamadas de *Validity Measures*, foram citadas por Babuska e Verbruggen (1997) e são merecedoras de melhor atenção.

Também merecem um estudo dedicado outras possíveis metodologias candidatas a trabalharem como sensores virtuais, como algoritmos de modelagem neuro-fuzzy e algoritmos nebulosos evolutivos. Este último, empregado com sucesso por Angelov e Filev (2002) inicialmente para uso em reconhecimento de padrões, é sem dúvida o mais sério e promissor dos candidatos.

BIBLIOGRAFIA CONSULTADA

ASCH, G. (ed.) **Les capteurs en instrumentation industrielle**. 5^e édition. Paris, Dunod, p.749-775, 1999.

AGUIRRE, L.A. **Introdução à identificação de sistemas: técnicas lineares e não-lineares aplicadas a sistemas reais**. 2.ed. Belo Horizonte: Editora UFMG, 2004. 659p.

ALLABOUTCIRCUITS. pH measurement. Disponível em <http://www.allaboutcircuits.com/vol_1/chpt_9/6.html>. Acesso em: 27 maio 2010.

ALVAREZ, H. et al. pH Neutralization process as a benchmark for testing nonlinear controllers. **Industrial & Engineering Chemistry Research**, vol. 40, n. 11, p.2467-2473, 2001.

ANGELOV, P.P. **Evolving rule-based models: a tool for design of flexible adaptive systems**. New York: Physica-Verlag Heidelberg, 2002. 213p.

ASCENCIO, R.R.L.; HERRERA, E. Sensores virtuales mediante redes neuronales artificiales. Dos estudios de caso en biotecnología. In: Memorias de la Conferencia Internacional CONIELECOMP, México. 1998. Disponível em: <<http://iteso.mx/~rleal/archivos/conielec.pdf>>. Acesso em: 9 mar. 2007.

ASCENCIO, R.R.L.; GALICIA, C.R.A. Biomass estimation using artificial neural networks on field programmable analog devices. In: The 2000 IEEE International Symposium on Industrial Electronics, Mexico. 2000. Disponível em: <<http://iteso.mx/~rleal/archivos/hardvirtual.pdf>>. Acesso em: 9 mar. 2007.

ATKINSON, C.M.; LONG, T.W.; HANZEVACK, E.L. Virtual Sensing: a neural network-based intelligent performance and emissions prediction system for on-board diagnostics and engine control. In: International Congress and Exposition, 1998, Detroit, Michigan, USA, **Anais...** Detroit: SAE, 1998, p.1-12.

BABUŠKA, R.; VERBRUGGEN, H.B. New approach to constructing fuzzy relational models from data. In: PROCEEDINGS THIRD EUROPEAN CONGRESS ON INTELLIGENT TECHNIQUES AND SOFT COMPUTING, 1995, Aachen, GERMANY, **Anais...** Aachen: EUFIT, 1995, p.583-587.

BABUŠKA, R.; VERBRUGGEN, H.B. Constructing fuzzy models by product space clustering. In: HELLENDORF, H.; DRIANKOV, D. (Ed) **Fuzzy model identification – selected approaches**. Berlin; Germany: Springer-Verlag Heidelberg, 1997. p.53-90.

BABUŠKA, R. **Fuzzy systems, modeling and identification**. 2000. 30 f. Knowledge-Based Control Systems (lecture notes for the course ET4-099), Delft University of Technology, Holanda. 2000. Disponível em: <<http://citeseerx.ist.psu.edu/viewdoc/summary?doi=10.1.1.29.9152>>. Acesso em: 20 abr. 2007.

BALDWIN, J.F.; MARTIN, T.P.; PILSWORTH, B.W. **FRIL – Fuzzy and evidential reasoning in artificial intelligence**. Taunton: Reaserch Studies Press, 1995. 406p.

BANERJEE, A. et al. H_{∞} control of non-linear processes using multiple linear models. In: PROCEEDINGS EUROPEAN CONTROL CONFERENCE, 1995, Rome, ITALY, **Anais...** Rome: ECC, 1995, p.2671-2676.

BERNARDO, L. D.; DANTAS, A.D.B. **Métodos e técnicas de tratamento de água**. 2ª edição. São Carlos, RIMA, 2005, 1566p.

BUPP, R. T.; BERNSTEIN, D. S.; COPPOLA, V. T. A benchmark problem for nonlinear control design: problem statement, experimental testbed, and passive nonlinear compensation. In: PROCEEDINGS OF THE AMERICAN CONTROL CONFERENCE, 1995, Seattle, USA, **Anais...** Seattle: ACC, 1995, p. 4363-4367.

CAMPOS, R. C. C. **Projeto e construção de uma planta piloto de neutralização de pH e proposta de metodologia para incorporação de informações auxiliares na identificação NARX racional**. 2007. 158p. Dissertação(Mestrado) – Centro Universitário do Leste de Minas Gerais, UnilesteMGol. Minas Gerais, 2007.

CASILLO, D. S. S. **Controle preditivo não linear baseado no modelo de Hammerstein com prova de estabilidade**. 2009. 135p. Tese(Doutorado) – Universidade Federal do Rio Grande do Norte, UFRN. Rio Grande do Norte, 2009.

CHEN, M.S.; WANG, S.W. Fuzzy clustering analysis for optimizing fuzzy membership functions. **Fuzzy Sets and Systems**, v.103, n.2, p.239–254, april. 1999.

COBBOLD, R. S. C. **Transducers for biomedical measurements: principles and applications**. New York: John Wiley & Sons, 1974, 486 p.

CYBOSOFT. GENERAL CYBERNATION GROUP. **MFA control and optimization on gas mixing process**. Sacramento, 2003. 3p. Disponível em: <http://www.cybosoft.com/pdf/ATS_15_GasMixingProcess.pdf>. Acesso em: 5 jun. 2007.

DEKHIL, H.; HENDERSON, T.C. Instrumented sensor systems architecture. **The International Journal of Robotics Research**, vol. 17, No. 4, p. 402-417, 1998.

DI BERNARDO, L.; DANTAS, A.D.B. **Métodos e técnicas de tratamento de água**. 5 ed. Rio de Janeiro: ABES, 1993, 1566p.

DIXON, K. et al. Rave: a real and virtual environment for multiple mobile robot systems. In: PROCEEDINGS OF THE IEEE INTERNATIONAL CONFERENCE ON ROBOTICS AND SYSTEMS, 1999, Kyongju, SOUTH KOREA, **Anais...** Kyongju: IROS, 1999, vol.3, p.1360-1367.

DORST, L. et al. Evaluating automatic debiting systems by modeling and simulation of virtual sensors. **IEEE Instrumentation and Measurement Magazine**, vol. 1, p.18-25, 1998.

ESMAILY-RADVAR, G. **Soft sensor development using nonlinear inferential modeling**. 2001. 8f. SEMINAR (CME DEPARTAMENTAL SEMINAR, CPC GROUP) - Department of Chemical and Materials Engineering, University of Alberta, Canada. 2001. Disponível em: <<http://www.ualberta.ca/CMENG/research/new-control/seminars>>. Acesso em: 9 mar. 2006.

FELTRE, R. **Química: físico-química**. vol.2. 5.ed. São Paulo: Editora Moderna, 2001, 560 p.

FORTESCUE, T.R.; KERSHENBAUM, L.S.; YDSTIE, B.E. Implementation of self-tuning regulators with variable forgetting factors. **Automatica**, v.17, n.6, p.831-835, 1981.

FORTUNA, L. et al. **Soft sensors for monitoring and control of industrial processes**. London: Springer-Verlag, 2007. 270p.

GONZALEZ, G.D. et al. Issues in soft-sensor applications in industrial plants. In: Symposium Proceedings, ISIE 94., IEEE Trans. Ind. Electron., 1994, p.380-385.

GONZALEZ, C.F.R.; ASCENCIO, R.R.L. Estimación de biomasa y pigmento en línea para una fermentación tipo fed-batch utilizando redes neurales artificiales. In: MEMORIAS DEL CONGRESO SOMI XV, 2000, Guadalajara, Jalisco, MEXICO, **Anais...** Guadalajara: SOMI, 2000, p. CIB-20-1.

GUSTAFSON, D. E.; KESSEL, W. C. Fuzzy clustering with a fuzzy covariance matrix. In: PROCEEDINGS OF IEEE CONFERENCE ON DECISION AND CONTROL, 1978, San Diego, USA, **Anais...** San Diego: CDC, 1978, p. 761-766.

GUSTAFSSON, T.K.; WALLER, K.V. Myths about pH and pH control. **Journal of American Institute of Chemical Engineers**, v. 32, n. 2, p. 335-337, 1986.

HALL, R. C; SEBORG, D. E. Modeling and self-tuning control of a multivariable pH neutralization process. In: PROCEEDINGS OF AMERICAN CONTROL CONFERENCE, 1989, Pittsburgh, USA, **Anais...** Pittsburgh: ACC, 1989, vol.2, p.1822-1827.

HENSON, M.; SEBORG, D. E. Adaptive nonlinear control of pH neutralization process. **IEEE Transaction on Control Systems Technology**, v.2, n.3, p.169–182, 1994.

HERNÁNDEZ, M.L.G; ASCENCIO, R.R.L.; GALICIA, C.A. An artificial neural network on a complex programmable logic device as a virtual sensor. In: PROCEEDINGS OF THE INTERNATIONAL SYMPOSIUM ON ROBOTICS AND AUTOMATION, 1998, Saltillo, Coahuila, MEXICO, **Anais...** Saltillo: ISRA, 1998, p.12-14.

JANG, J.S.R. Self-learning fuzzy controllers based on temporal back propagation. **IEEE Transaction on Neural Networks**, v.3, n.5, p.714–723, 1992.

JERÔNIMO, R. A. **Seleção de variáveis de entrada para identificação de sistemas aplicada a uma planta de tratamento de efluentes**. 2004. 280 p. Tese (Doutorado) - ESCOLA POLITECNICA, Universidade de São Paulo, São Paulo, 2004.

JIN, Y.; JIANG, J.; ZHU, J. Adaptive fuzzy modeling and identification with its applications. **International Journal of Systems Science**, v.26, n.2, p.197–212, 1995a.

JIN, Y.; JIANG, J.; ZHU, J. Neural network based fuzzy identification and its application to modeling and control of complex systems. **IEEE Transaction on Systems, Man and Cybernetics**, v.25, n.6, p.990–997, 1995b.

JOHANSEN, T.A.; FOSS, B.A. Constructing narmax models using armax models. **Intenational Journal of Control**, v.58, p.1125–1153, 1993.

KAYMAK, U.; BABUŠKA, R. Compatible cluster merging for fuzzy modeling. In: PROCEEDINGS OF IEEE INTERNATIONAL CONFERENCE ON FUZZY SYSTEMS, 1995, Yokohama, JAPAN, **Anais...** Yokohama: FUZZ-IEEE/IFES, 1995, p. 897 –904.

KESTELL, C.D. **Active control of sound in a small single engine aircraft cabin with virtual error sensors**. 2000. 255p. Tese (Doutorado) - Department of Mechanical Engineering, Adelaide University. Australia, 2000. Disponível em: <<http://thesis.library.adelaide.edu.au/adt-SUA/uploads/approved/adt-SUA20010216.164243/public/02whole.pdf>>. Acesso em : 9 mar. 2007.

LJUNG, L; TORDEL, G. **Modeling of dynamic systems**. 1st ed. New Jersey: Prentice-Hall, 1994. 368p.

LJUNG, L. **System identification: theory for the user**. 2nd ed. New Jersey: Prentice-Hall, 1999. 603p.

LUCENA, S. E. de. The electronic detail of a digital pH-meter. In: IMEKO XVIII World Congress (International Measurement Confederation), 2006, Rio de Janeiro, Brasil, **Anais...** Rio de Janeiro: IMEKO, 2006. v. CD-ROM.

McMILLAN. **pH Measurement and control**. 2nd ed. Research Triangle Parc: ISA, 1994. 294p.

MENDEL, J.M.; Fuzzy logic systems for engineering: a tutorial. **Proceedings of the IEEE**, v.83, n.3, p.345-377, 1995.

MERRITT, R. Hardware vs. virtual analyzers. **Control Magazine**. Março, 2002. Aplicações de sensores virtuais em fabricação de papel. Disponível em: <http://www.controlmagazine.com/Web_First/ct.nsf/ArticleID/PSTR-56KM99>. Acesso em: 06 set. 2007.

MOUSTAFA, K.A.F. Identification of stochastic time-varying systems. **Proceedings of the Institute of Electrical Engineers**, v.130, Pt D, n.4, p.137-142, 1983.

MUIR, P.F. A virtual sensor approach to robot kinematic identification: theory and experimental implementation. In: PROCEEDINGS OF THE IEEE INTERNATIONAL CONFERENCE ON SYSTEMS ENGINEERING, 1990, Pittsburgh, USA, **Anais...** Pittsburgh: ICSYSE, 1990, p. 440-445.

NARENDRA, K. S.; PARTHASARATHY, K. Identification and control of dynamical systems using neural networks. **IEEE Transaction on Neural Networks**. v. 1, p.4-27, 1990.

NATIONAL INSTRUMENTS CORP. **pH measurement tutorial**. Disponível em: <<http://zone.ni.com/devzone/cda/tut/p/id/2870> >. Acesso em 27/05/2010.

NGUYEN, H. T.; SUGENO, M. (Ed) **Fuzzy systems: modeling and control**. Boston, USA: Kluwer Academic Publishers, 1998.

NIE, J.; LOH, A. P.; HANG, C. C. Modeling pH neutralization processes using fuzzy-neural approaches. **Fuzzy Sets and Systems**, v.78, p. 5–22, 1996.

OLIVEIRA, C.E.N.; GARCIA, C. Aplicação de lógica fuzzy no desenvolvimento de um sensor virtual. Relatório apresentado ao FDTE – Fundação para o Desenvolvimento Tecnológico da Engenharia. São Paulo, fev. 2003.

OLIVEIRA, J.V. Neuron inspired rules for fuzzy relational structures. **Fuzzy Sets and Systems**, v.57, n.1, p.41–55, 1993.

OOSTEROM, M.; BABUŠKA, R. Virtual sensor for fault detection and isolation in flight control systems – fuzzy modeling approach. Proceedings of 39th IEEE Conference on Decision and Control, Sydney, December 2000. Disponível em: <http://dutera.et.tudelft.nl/~oosterom/papers/2000/cdc00_inv5806.pdf>. Acesso em: 9 mar. 2003.

PAIVA, S.; MANZI, J. Modelling, identification and simulation applied to neutralization system. **Internation Journal of Modelling, Identification and Control**, v. 2, n. 3, p. 235–240, 2007.

PEDRYCZ, W. Relevancy of fuzzy models. **Information Sciences**, v.52, p. 285–302, 1990.

ROSS, T.J. **Fuzzy logic with engineering applications**. 2nd Ed., Chichester: John Wiley & Sons Ltd., 2004. 650p.

SABADIN, V. A evolução da medição de pH pelos eletrodos digitais. **InTech America do Sul – ISA Distrito 4**, n. 105, p. 24-33, 2008.

SAFAVI, A.; BAGHERI, M. Novel optical pH sensor for high and low pH values. **Sensors and Actuators**, B, 90, p.143-150, 2003.

SANTOS, J. E. S. **Controle preditivo não linear para sistemas de Hammerstein**. 2007. 156p. Tese (Doutorado) – Universidade Federal de Santa Catarina. Florianópolis, 2007.

SCHANZER, S. **Identificação de sistemas utilizando modelos nebulosos**. 2002. Dissertação (Mestrado) – Escola Politécnica, Universidade de São Paulo. São Paulo, 2002.

SEBORG, D. E.; EDGAR, T. F.; MELLICHAMP, D. A. **Process dynamics and control**. 2nd Ed., Hoboken: John Wiley & Sons Ltd., 2004. 697p.

SHARMA, N.; GUPTA, B.D. Fabrication and characterization of pH sensor based on side polished single mode optical fiber. **Optics Communications**, n.216, p.299-303, 2003.

SHINSKEY, G. **Process control systems: application, design, and tuning**. 4nd Ed., Boston: McGraw Hill, 1996. 435p.

SIEGEL, M. Sensor modeling and simulation: Can it pass the Turing Test? IEEE International Workshop on Virtual and Intelligent Measurement Systems. Budapest, May 19-20, 2001.

SOFT SENSOR – WASTEWATER, Pavilion Technologies. Aplicações de soft-sensor em tratamento de água. Disponível em: <[http://www.pavtech.com/pavilion /value/industry/environmental/Environmental_Water.html](http://www.pavtech.com/pavilion/value/industry/environmental/Environmental_Water.html)>. Acesso em: 06 de set. 2007.

SOTOMAYOR, O. A. Z. **Modelamento, simulação e controle de um processo de neutralização de pH**. 1997. 165 p. Dissertação (Mestrado) – Escola Politécnica, Universidade de São Paulo. São Paulo, 1997.

SUGENO, M; TANAKA, K. Successive identification of a fuzzy model and its application to prediction of a complex systems. **Fuzzy Sets and Systems**, v.42, p.315–324, 1991.

TAKAGI, T.; SUGENO, M. Fuzzy identification of systems and its application to modeling and control. **IEEE Trans. Systems, Man and Cybernetics**, v.15, n.1, p.116–132, 1985.

TANAKA, K.; WANG, H.O. **Fuzzy control systems design and analysis: a linear matrix inequality approach**. New York: John Wiley & Sons, 2001.

US BIOSOLUTIONS BRASIL: soluções em equipamentos e serviços. Disponível em: <<http://www.usbio.com.br>>. Acesso em: 23 de out. 2010.

VAN GORP, J.; SCHOUKENS, J. A scheme for nonlinear modeling. In: 5th International Conference on Information Systems Analysis and Synthesis, ISAS'99, Orlando, USA, 1999, p.450-456.

VOUTCHKOV, I. I.; VELEV, K. D. Identification of waste water treatment plant using neural networks. In: REUSCH, B. (Ed) **Computational intelligence theory and application – Lecture notes in computer science**. Dortmund; Germany: Springer-Verlag, vol. 1226. 1997. p. 478-483.

WALLER, K. V.; MÄKILÄ, P. M. Chemical reaction invariants and variants and their use in reactor modelling. **Industrial and Engineering Chemistry, Process Design and Development**, v.38, p.1-11, 1981.

WANG, L.-X. **Adaptive fuzzy systems and control, design and stability analysis**. New Jersey: Prentice-Hall, 1994.

WANG, L.; LANGARI, R. Identification of time-varying fuzzy systems. In: Proceedings of the First International Joint Conference of the North American Fuzzy Information Processing Society Biannual Conference, 1994. NAFIPS/IFIS/NASA '94: The Industrial Fuzzy Control and Intelligent Systems Conference, and the NASA Joint Technology. p.365-369, Dec. 1994.

WANG, L.; LANGARI, R. Complex systems modeling via fuzzy logic. **IEEE Transaction on Systems, Man and Cybernetics – Part B: Cybernetics**, v.26, n.1, p.100-106, Feb. 1996a.

WANG, L.; LANGARI, R. A variable forgetting factor RLS algorithm with application to fuzzy time-varying systems identification. **International Journal of Systems Science**, v.27, n.2, p.205-214, 1996b.

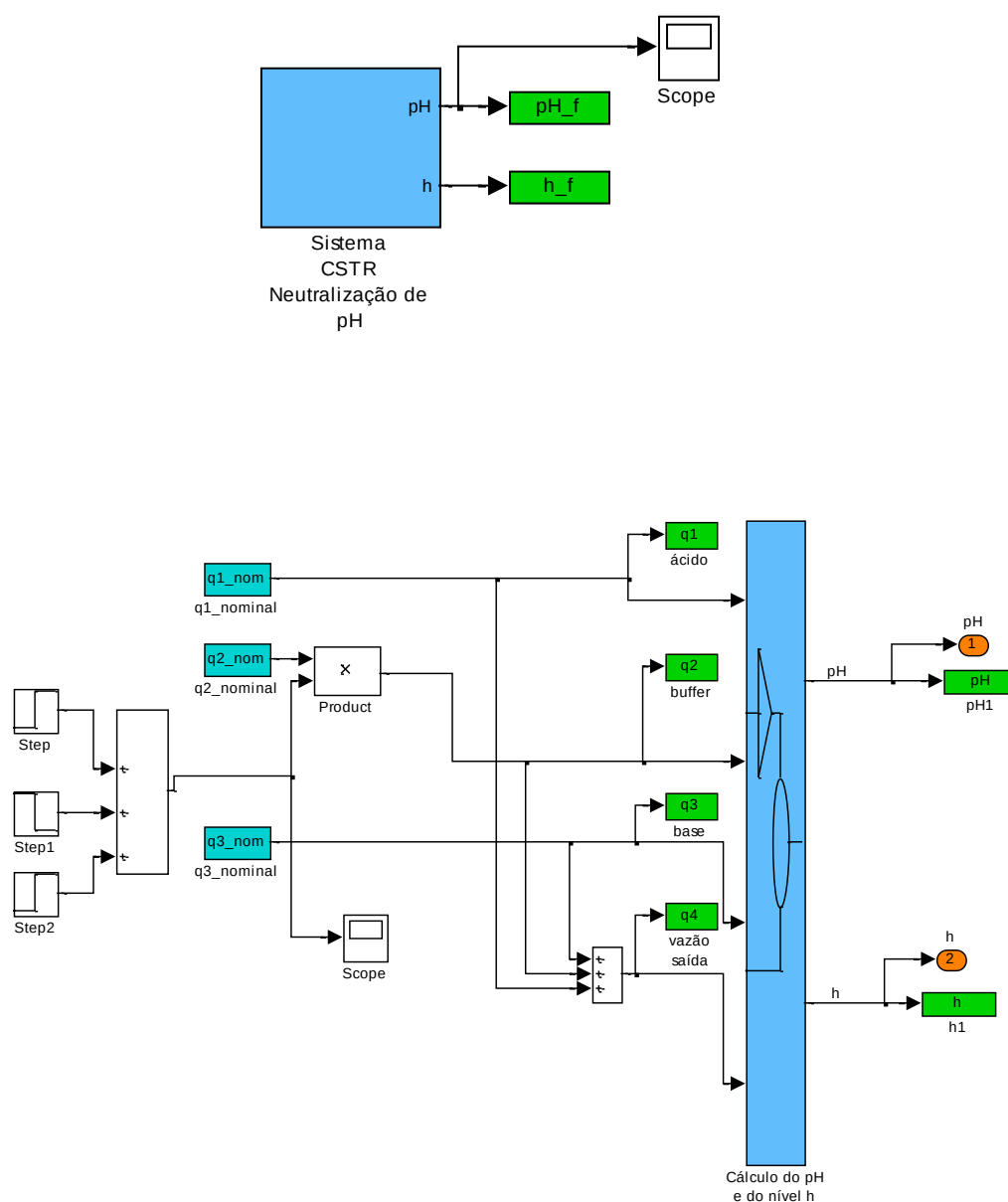
YEN, J.; Fuzzy logic – A modern perspective. **IEEE Transactions on Knowledge and Data Engineering**, v.11, n.1, p.153-165, 1999.

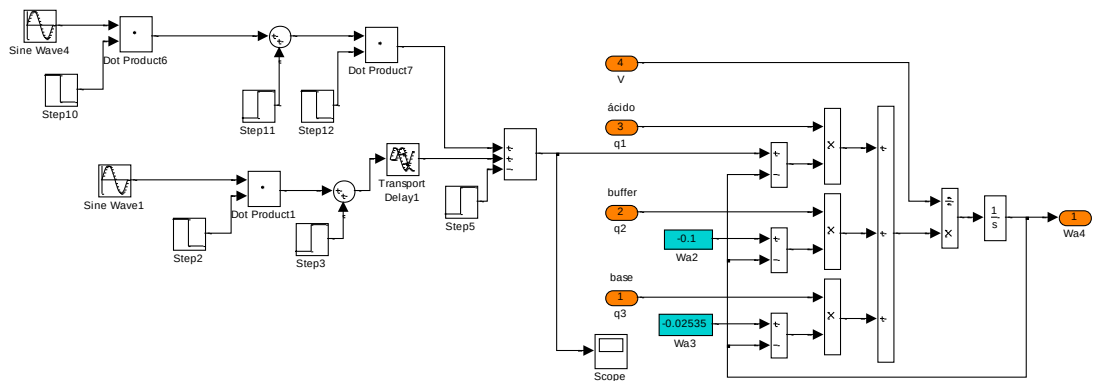
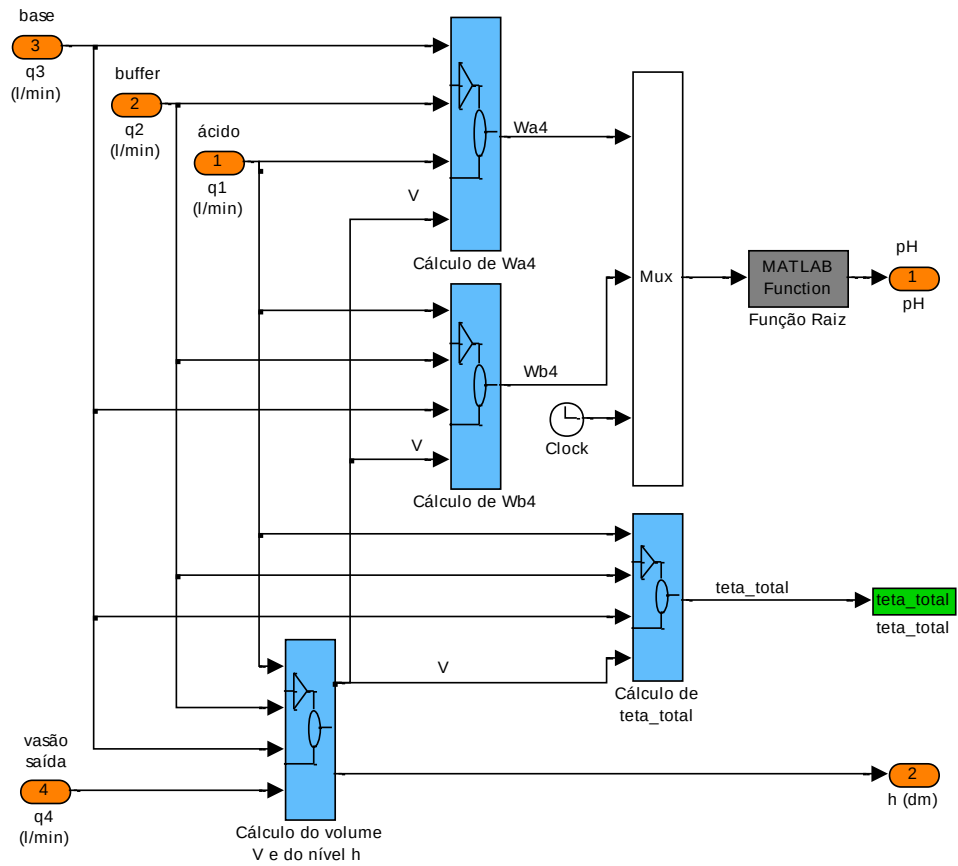
ZADEH, L.A.; Fuzzy sets. **Information and Control**, v.8, p.338-352, 1965.

ZADEH, L.A.; Outline of a new approach to the analysis of complex systems and decision processes. **IEEE Transactions on Systems, Man and Cybernetics**, v. SMC-3, n.1, p.28-44, 1973.

APÊNDICE A – Modelo virtual do processo de neutralização de pH feito em Matlab®/Simulink®

O modelo descrito abaixo representa a modelagem do processo de neutralização e mais especificamente, com os ajustes e condições necessários para simulação do teste 4 apresentado no Capítulo 4.





Função raiz

```
function y=raiz(u);

wa4=u(1); wb4=u(2); t=u(3);

global pHi

ka1=4.47e-7;
ka2=5.62e-11;
kw=1e-14;
teta=10/60;    %time delay na medição do pH

n1=-1;
n2=wa4-ka1;
n3=ka1*(wb4+wa4-ka2)+kw;
n4=2*wa4*ka1*ka2+4*wb4*ka1*ka2+ka1*kw; %n4=-ka1*(kw+ka2*wa4+2*ka2*wb4);
n5=ka1*ka2*kw;

format long

va=roots([n1 n2 n3 n4 n5]);

for cont=1:4
    if imag(va(cont))==0 & real(va(cont))>0
        pH=-log10(va(cont));
    end
end

y=pH;
```

Dados de inicialização

```
% Dados do CSTR

A=6.7887;          % Área do CSTR
h_nom=2.94609;     % Nível nominal no CSTR

delta_f=(2*1000*(10/60)+20); delta=0.0:(10/60):delta_f;
Ni=-(100/delta_f)*delta+100;          % Velocidade da hélice do agitador [rpm]
Nq=0.4;          % Coeficiente de descarga do agitador
Dt=0.2942/0.3048;          % Diâmetro interno do CSTR
Di=0.3*Dt;          % Diâmetro do círculo formado pela hélice do agitador
Fa=(7.4813*Ni*Nq*3.7854*Di^3)/(0.3^0.55);' % Vazao de fluido movimentada pelo agitador

% Tempo morto do pH e intervalo de amostragem
teta_pH=10/60;          % Tempo morto do pH [min]
Ts=10/60;          % Intervalo de amostragem [min]

% Condições iniciais
Wa40=-0.0018862; Wb40=0.0023067;
V0= A*h_nom;          % Volume inicial no CSTR

% Condições nominais
q1_nom=0,9996;          % Vazão nominal de ácido
q2_nom=0,0333;          % Vazão nominal de buffer
q3_nom=0.4686;          % Vazão nominal de base
q4_nom=q1_nom+q2_nom+q3_nom;          % Vazão nominal de saída
C1_nom=2;          % Concentração nominal do ácido
```

APÊNDICE B – *Script-file* em Matlab® para aplicação ao primeiro exemplo do item 5.1

```
%WANG,L.; LANGARI, R. A variable forgetting factor RLS algorithm with application to fuzzy time-
%varying systems identification. International Journal of Systems Science, v. 27, nr. 2, pp. 205-214, 1996.

%Slow case
clear all

n=1000;      %Número de dados
clusters=6;  %Número de clusters
num=2;      %Número de variáveis de autoregressão

y(1)=0;      %Saída inicial
a(1)=.8;
b(1)=2.0;
u(1)=.7*sin(1);
dado=[y(1) u(1)];

for k=2:n
    if k<=500
        a(k)=(1-k^(-2))*a(k-1);
        b(k)=(1-k^(-2))*b(k-1);
        y(k)=a(k)*y(k-1) + b(k)*u(k-1);
    end;
    if k>500
        a(k)=(1-k^(-1))*a(k-1);
        b(k)=(1-k^(-1))*b(k-1);
        y(k)=a(k)*y(k-1) + b(k)*u(k-1);
    end;
    dado=[dado;[y(k-1) u(k-1)]];
    u(k)=.7*sin(k);
end;

%Processo de clusterização de Gustafsson-Kessel
%BABUSKA, R.; VERBRUGGEN, H.B. Constructing fuzzy models by product space clustering. Fuzzy Model
%Identification: Selected Approaches, pp. 53-90, 1997

Z=dado;
U0=rand(n,clusters);
m=2;
tol=5e-4;

[U,V,F]=guskel(Z,U0,m,tol);

%Montagem da matriz Phi

Phi=[];
for i=1:n
    %Phi=[Phi;[U(i,:) U(i,:)*Z(i,1) U(i,:)*Z(i,2)]];
    Phi=[Phi;[U(i,:)*Z(i,1) U(i,:)*Z(i,2)]]; %No caso do exemplo, não há termo independente
end;

%Montagem do Variable Forgetting Factor RLS Algorithm

%Montagem da matriz Ck inicial (triangular inferior)
aux=size(Phi,2);
Mat=ones(aux);
```

```

Ck=tril(Mat);

%Constantes iniciais de convergência
lambda(1)=1;
lambdamin=.2;
lambdamax=1;
tetak=zeros(aux,1);
delta=.99;
tau=1000;

%Montagem do looping principal
for k=1:n
    fk=Ck'*Phi(k,:);
    gamak=(lambda(k)+fk'*fk)^.5;
    alfak=gamak+lambda(k)^.5;
    betak=1/(gamak*alfak);
    sigmak=betak*Ck*fk;
    e(k)=y(k)-Phi(k,:)*tetak;
    tetak=tetak+sigmak*(alfak/gamak)*e(k);
    lambda(k+1)=delta-tau*e(k)^2;
    if lambda(k+1) < lambdamin
        lambda(k+1)=lambdamin;
    end;
    if lambda(k+1) > lambdamax
        lambda(k+1)=lambdamax;
    end;
    Ck=lambda(k+1)^(-.5)*(Ck-sigmak*fk');
    a1(k)=U(k,:)*tetak(1:aux/2);
    b1(k)=U(k,:)*tetak(aux/2+1:aux);
end;

% Entrada
figure(1)
plot(1:n,u); grid; titulo2 = [' Sinal de entrada ']; title(titulo2); ylabel(' uk '); xlabel('amostras');

% Saida
figure(2)
plot(1:n,y); grid; titulo2 = [' Sinal de Saída ']; title(titulo2); ylabel(' yk '); xlabel('amostras');

% Parametro a
figure(3)
plot(1:n,a,1:n,a1,'r:'); titulo2 = [' Parâmetro a real e estimado ']; title(titulo2); ylabel(' Parâmetro a ');
xlabel('amostras');

% Parametro b
figure(4)
plot(1:n,b,1:n,b1,'r:'); titulo2 = [' Parâmetro b real e estimado ']; title(titulo2); ylabel(' Parâmetro b ');
xlabel('amostras');

% Evolução do fator de esquecimento
figure(5)
plot(1:n,lambda(1:n)); grid; titulo6 = [' Evolução do fator de esquecimento variável '];
title(titulo6); ylabel(' lambda '); xlabel('amostras')

% Evolução do valor RMS do Erro de predição
figure(6)
plot(1:n,e); grid; titulo3 = [' Evolução do Erro de predição ']; title(titulo3); ylabel(' Erro saída ');
xlabel('amostras');

```