

## Sistema de Inferência da Fração Molar dos Contaminantes do Gás Liquefeito de Petróleo Baseado em Técnicas de Inteligência Artificial

Jean Mário Moreira de Lima, jean@dca.ufrn.br<sup>1</sup>

Fábio Meneghetti Ugolino de Araújo, meneghet@dca.ufrn.br<sup>1</sup>

<sup>1</sup>UFRN/CT, DCA - Laboratório de Controle de Processos, Lagoa Nova, Natal/RN - Brasil, CEP 59.078-970

**Resumo:** Em UPGNs (Unidades de Processamento do Gás Natural), um dos produtos de maior rentabilidade é o GLP (Gás Liquefeito de Petróleo) que é composto majoritariamente por propano (C3) e butano (C4). Além disso, o pentano (C5) e o etano (C2) também podem se fazer presente no gás como contaminantes. A medição da fração molar dos componentes do GLP é feita através de cromatógrafos a gás. Porém a cromatografia é um processo lento, impossibilitando que o monitoramento da qualidade do produto seja feito em tempo real. Neste trabalho, propõe-se um sistema de inferência, utilizando a redes neurais artificiais, com o objetivo de inferir as frações molares do C5 e do C2. Dessa forma, seria possível estimar às frações molares desses contaminantes do GLP minuto a minuto, de maneira consideravelmente mais rápida que a tradicional por cromatógrafos. Os resultados obtidos são promissores, mostrando que o sistema de inferência pode inferir os frações molares de C5 e C2 e, assim, habilitar melhora no monitoramento da qualidade do GLP e, consequentemente, na lucratividade.

**Palavras-chave:** UPGN, GLP, Sistema de Inferência, Inteligência Artificial, Redes Neurais Artificiais

### 1. INTRODUÇÃO

Atualmente, diante de um mercado cada vez mais competitivo, produzir de forma eficiente e eficaz é essencial para se obter um balanço econômico positivo. Reduzir custos, realizar processos otimizados e ofertar produtos cada vez melhores são fatores que influenciam diretamente na economia de qualquer indústria. Diante disso, técnicas que podem melhorar e/ou garantir otimização de processos, como o monitoramento da qualidade do produto ou de controle avançado e inteligente tornam-se fundamentais para a indústria como um todo.

No caso de Unidades de Processamento de Gás Natural (UPGNs), o monitoramento da qualidade do produto produzido é intrínseco a uma produção satisfatória, e esse controle da qualidade faz-se, como na maioria dos processos químicos, através da composição química dos produtos. Composições químicas são variáveis de difícil mensuração (Linhares, 2010). De acordo com Warne et al. (2004), a medição da composição química de um componente é normalmente feita através de análise em laboratório. Dessa forma, longos intervalos de tempo para obtenção das amostras e análise laboratorial são requeridos.

De acordo com Bolf et al. (2008), instrumentos de medição online da composição química de produtos até estão disponíveis como os cromatógrafos à gás e analisadores NIR (Near-InfraRed), mas apresentam custo elevado e longos intervalos de medição.

Quando limitações técnicas, custo elevado e indisponibilidade impedem o uso de sensores de variáveis primárias em tempo real, Zamprogna et al (2005) afirma que sistemas de inferência apresentam uma forma atraente de se lidar com o problema de medição dessas variáveis de processo. De acordo com Ivandic et al (2008), o objetivo de desenvolver um sistema de inferência é modelar uma relação entre as variáveis primárias, de difícil medição, com as variáveis ditas secundárias, de fácil e rápida medição.

Abdalla et al (2008) apresentou um sistema de inferência para predição da composição molar do pentano normal (n-pentano) e o isopentano (i-pentano) em uma coluna de destilação debutanizadora. Utilizando-se redes neurais artificiais, implantou-se um sistema de inferência que recebe como entrada as variáveis primárias de fácil medição: temperatura, vazão de refluxo a vazão, e têm como saídas: n-pentano e i-pentano. Os resultados mostraram-se satisfatórios, com coeficiente de correlação entre valores reais e estimados igual a 0.999.

Ainda na mesma linha de inferência, Hongmei et al. (2015) implementou um *soft sensor* que pudesse fazer a inferência da concentração de butano no fundo de uma coluna de destilação debutanizadora. Para desenvolver este *soft sensor*, foi proposto um método baseado em regressão ortogonal dinâmica. O método é aplicado a variáveis secundárias atrasadas em uma ordem pré-determinada. As variáveis que apresentam maior relação com o comportamento das variáveis primárias são selecionadas para construir o modelo do *soft sensor*. As simulações mostraram a eficiência do método utilizado e da aplicação de sensoriamento virtual na estimação da concentração de compostos químicos em uma coluna de destilação utilizando *soft sensor*.

Normalmente, o produto de maior rentabilidade econômica em uma UPGN é o GLP. De forma ideal, é formado por propano e butano (C3 e C4). Na prática, porém, há contaminantes como o etano e o pentano (C2 e C5). O sistema de inferência aqui proposto tem como principal objetivo estimar as frações molares dos principais contaminantes que constituem o GLP (variáveis primárias), isto é, etano e pentano a partir de várias de fácil medição como temperatura e pressão (variáveis secundárias). A ideia é que esse sensor virtual possa fazer a função que um cromatógrafo a gás faz, gerando percentuais das composições dos contaminantes do GLP, só que em tempo real. Dessa forma, o monitoramento da qualidade do GLP produzido torna-se possível, uma vez que a medição da composição do GLP não será através do lento processo analítico. Assim, melhora-se a qualidade do processo, do produto e, conseqüentemente, sua lucratividade.

Para implementação do soft sensor deste trabalho, aplica-se a técnica de inteligência artificial conhecida como redes neurais artificiais. De acordo com Linhares et al (2010), são poderosas ferramentas para a identificação de dinâmicas não lineares complexas, como em uma UPGN. No escopo deste trabalho, o principal benefício é que elas podem associar as relações dinâmicas não-lineares entre as variáveis primárias e secundárias de um processo.

Para o desenvolvimento deste trabalho, utilizou-se uma UPGN simulada em software HYSYS, baseada na UPGNII GMR, da Petrobrás, que fica na cidade de Guamaré/RN. É formada por uma coluna deetanizadora em série com uma coluna debutanizadora. Na instrumentação da planta, têm-se alguns controladores PID's. O sistema de inferência é baseado a partir de variáveis do processo desses controladores. Neste trabalho, é proposto também um sistema de correção do erro do sistema de inferência, tendo com base a leitura da composição do GLP feita por cromatógrafos presentes no processo.

## 2. FUNDAMENTAÇÃO TEÓRICA

### 2.1 Redes Neurais Artificiais

De acordo com Haykin (2001), uma rede neural artificial pode ser definida como uma estrutura maciçamente paralela, distribuída e formada em sua totalidade por unidades simples de processamento: os neurônios. São capazes de aprender a partir do ambiente através de um processo de aprendizagem e esse conhecimento é armazenado pela força de conexão entre os neurônios. Ainda de acordo com Haykin (2001), as principais propriedades de uma rede neural artificial são: capacidade de generalização, adaptabilidade, não-linearidade intrínseca, mapeamento entrada-saída.

Linhares et al. (2008) afirma que através de características como a capacidade de aproximação universal e aprendizagem, uma rede neural artificial pode representar a dinâmica de um sistema complexo de forma relativamente simples. Reforça-se, assim, o conceito de utilização de redes neurais em aplicações de comportamento dinâmico. Essa capacidade das redes é vital a este trabalho, já que se deseja inferir o comportamento de um sistema complexo.

Há diversas arquiteturas de redes neurais artificiais. A forma como os neurônios da rede estão estruturados é intrinsecamente ligada ao algoritmo de aprendizado utilizado no treinamento da rede. Uma vez que apresentou resultados satisfatórios em sistemas de inferências, como de Fortuna et al. (2005), a arquitetura utilizada neste trabalho será a perceptron de múltiplas camadas (PMC).

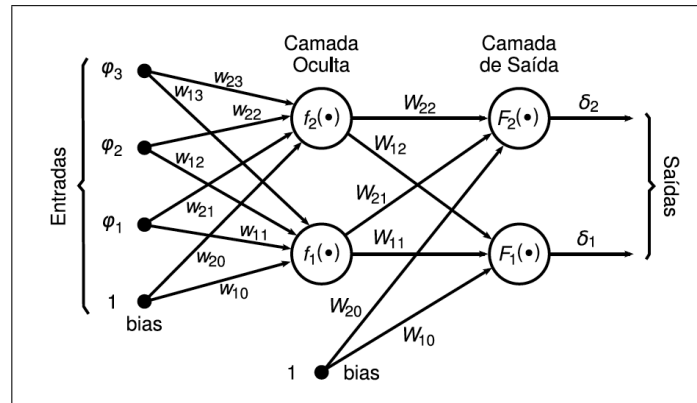


Figura 1: PML interconectada de duas camadas com dois neurônios.

Uma rede PMC, também chamada de *feedforward*, é constituída basicamente por uma camada de entrada, uma ou mais camadas ocultas nas quais os neurônios estão interconectados e uma camada de saída. Nesse esquema, cada neurônio recebe o sinal da camada imediatamente anterior ou das entradas. Este tipo de rede é totalmente interconectada, ou seja, todas as entradas e neurônios de uma camada estão conectados com todos os neurônios da camada seguinte. É possível sintetizar o funcionamento de uma rede PMC matematicamente:

$$\delta_i(t) = g_i(\varphi, \theta) = F_i \left[ \sum_{j=1}^{n_h} W_{i,j} f_j \left( \sum_{l=1}^{n_\varphi} w_{j,l} \varphi_l + w_{j,0} \right) + W_{i,0} \right] \quad (1)$$

Sendo  $\theta$  um vetor dos parâmetros ajustáveis: biases e pesos sinápticos ( $W_{i,j}, w_{j,0}$ ).  $n_h$  e  $n_\varphi$  são respectivamente o número de neurônios na camada oculta e o número de entradas da rede. Por fim, tem-se  $f_i$  e  $F_i$  que são as funções de

ativação dos neurônios nas camadas oculta e de saída, respectivamente. Comumente, as funções de ativação utilizadas são a linear, sigmóide e atangente hiperbólica

### 2.1.1 Aprendizagem

De acordo com Haykin (2001), uma das mais importantes propriedades de uma rede neural artificial é sua capacidade de aprender a partir de determinado ambiente e consequentemente melhorar sua performance com aprendizagem. Esse aprendizado deve-se aos ajustes dos parâmetros ajustáveis biases e pesos sinápticos a cada iteração do processo de aprendizagem.

O processo de aprendizagem supervisionado é o principal método de aprendizagem aplicado às redes PMC. A principal característica desse método é quando há disponibilidade de dados do sistema ou ambiente na forma de padrões de entrada-saída. Isto é, para uma certa entrada, a saída da rede neural é comparada com a saída correspondente do conjunto de treinamento. Assim, baseados no sinal de erro e nos padrões de treinamento, os parâmetros livres da rede podem ser ajustados.

As regras estabelecidas para lidar com esse problema de aprendizagem, isto é, o ajuste citado acima, é denominado algoritmo de aprendizagem ou de treinamento. Há diversos algoritmos de treinamentos de redes neurais, e a diferenciação entre eles se dá pela forma como os parâmetros adaptáveis (biases e pesos sinápticos) de uma rede são ajustados. O algoritmo que apresenta-se como o mais usual no processo de aprendizagem supervisionado de redes PML é o de retropropagação de erros ou, do inglês, backpropagation. Por sua vez, é o algoritmo escolhido para implementação deste trabalho.

## 2.2 Identificação de Sistemas

Inferir o comportamento de um sistema, é representar a dinâmica do mesmo de forma satisfatória por alguma técnica. De acordo com Aguirre (2004), há várias técnicas de modelar a dinâmica de um sistema, identificando-o. Basicamente, deve-se considerar o conhecimento que se tem sobre o sistema que se deseja inferir. Assim, classifica-se a técnicas em: caixa branca, caixa preta e caixa cinza. Caso a identificação baseie-se apenas em dados coletados do sistema e assumindo que nenhum ou pouquíssimo conhecimento sobre o sistema é usado, a técnica de modelagem é dita caixa preta.

A realização da identificação de um sistema caixa preta assemelha-se as etapas do projeto de redes neurais. Logo, devido a essa semelhança e a características intrínsecas das redes, a identificação de sistemas complexos tem sido feita por RNAs que conseguem representar a dinâmica de forma satisfatória e relativamente simples. Neste trabalho, utiliza-se a modelagem caixa preta uma vez que o sistema que se deseja identificar é complexo e não há conhecimento de sua modelagem fenomenológica. Coleta-se dados experimentais do sistema simulado formado pelas colunas de destilação deetanizadora e debutanizadora com o objetivo de obter o sistema de inferência baseado em um modelo neural capaz de representar as dinâmicas entre as variáveis primárias e secundárias do sistema.

### 2.2.1 Identificação Utilizando Modelos Neurais

Basicamente, as estruturas de modelagem com base em RNAs com características para identificar sistemas não lineares, são generalizações das que apresentam modelagem linear. São exemplos de estruturas de modelagem linear: modelo ARX (AutoRegressive eXogenous input), ARMAX (AutoRegressivo, Moving Average, eXogenous input), OE (Output Error). Esses modelos são caracterizados por um vetor que armazena valores passados das variáveis utilizadas na estimativa da saída. Esse vetor chama-se vetor de regressão e, de acordo com a escolha desse vetor, diferentes estruturas de modelos neurais podem ser obtidos. Por exemplo, se o vetor de regressão escolhido é equivalente ao utilizado no modelo ARX, a estrutura de modelo neural é chamada NNARX (Neural Network ARX).

No modelo NNARX, o vetor de regressão, além de armazenar os valores passados das variáveis utilizadas na estimativa da saída, também guarda os valores passados das variáveis de saída do processo. De acordo com Norgaard (2001), esse modelo é estável uma vez que apresenta relações puramente algébricas entre suas variáveis. Devido a isso, a estrutura de modelagem neural escolhida como base a ser utilizada no sensor virtual aqui proposto é o modelo NNARX ilustrada na Figura 2.

A Figura 2 apresenta o diagrama de uma estrutura de modelagem NNARX. Nesta estrutura  $\hat{y}$  é a saída da rede, isto é, a estimativa da saída da planta.  $a$  é o atraso de transporte da planta,  $n$  é a ordem de entrada da planta,  $m$  é a ordem de saída da planta,  $y$  é a saída da planta e  $u$  a entrada. Também é possível observar regressores. Sabendo-se que o projeto da estrutura é feito para identificar a dinâmica de um sistema físico, a utilização de regressores é intrínseca pois os mesmos fazem com que a saída da rede neural esteja diretamente relacionada com valores anteriores tanto de entrada como de saída do sistema.

O modelo NNARX exposto será utilizado para treinamento do sistema de inferência. Entretanto, como descrito acima, nas suas entradas, também estará presente as saídas passadas do sistema, ou seja, variáveis primárias de difícil medição. Sendo assim, o sistema de inferência não seria capaz de reduzir o tempo de medição dos cromatógrafos, já que seriam necessárias informações das medições dos mesmos. Para contornar isso, após treinamento, utilizaremos como entrada, para os testes e validações posteriores, os valores passados estimados pelo próprio modelo neural. Assim, garante-se que a avaliação do modelo seja com relação ao seu potencial de inferir as variáveis primárias a partir das secundárias e de suas próprias estimativas em tempo de execução.

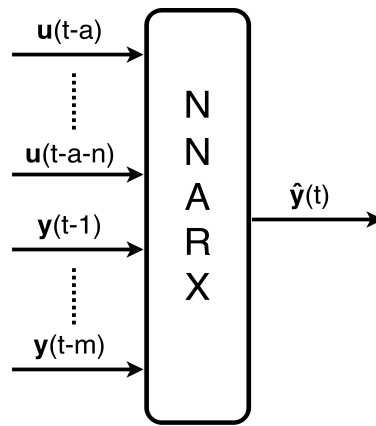


Figura 2: Estrutura de modelagem NNARX.

### 3. METODOLOGIA

#### 3.1 Módulo RNA

##### 3.1.1 Saídas

A inferência de variáveis primárias de difícil medição a partir de variáveis secundárias de fácil medição para monitoramento é o objetivo deste trabalho. Por isso, adota-se a fração molar dos principais contaminantes do GLP como as variáveis primárias, uma vez que a medição das mesmas através de cromatógrafos a gás é lenta. Sendo assim, as saídas do modelo neural proposto são ilustradas na Tabela 1:

Tabela 1: Variáveis primárias do processo.

| p | Variável Primária ( $VP_i$ ) |
|---|------------------------------|
| 1 | Fração Molar do Etano (C2)   |
| 2 | Fração Molar do Pentano (C5) |

##### 3.1.2 Entradas

A seleção das entradas, ou variáveis secundárias que servirão de instrumento para a inferência das variáveis de processo, isto é, as variáveis primárias, é importante uma vez que a quantidade dessas variáveis vai influenciar no grau de qualidade e na complexidade do sistema de inferência. De acordo com Linhares et al 2010, Em sistemas de coluna de destilação, comumente, podem ser variáveis secundárias temperaturas dos diferentes pratos perfurados, vazão de refluxo, pressão de coluna, temperatura do refeedor e condensador e volume de condensado.

Como o sistema simulado buscar manter características reais da planta UPGN-II GMR, apenas leva-se em consideração os valores de variáveis secundárias que podem ser obtidos através de sensores físicos reais. A motivação é que deve se proposto um sistema de inferência eficiente sem que haja necessidade de instalar novos instrumentos na planta.

Para a inferência de pentano e do etano, escolheu-se a temperatura do estágio 16 e vazão de refluxo da coluna debutanizadora e temperatura do estágio 40 e o percentual de volume de condensado da coluna deetanizadora, respectivamente. A escolha é feita com base na influência dinâmica que as variáveis secundárias impõe sobre as primárias, isto é, a correlação entre as mesmas.

Tabela 2: Variáveis secundárias do processo.

| p | Variável Secundária ( $VS_p$ ) | Controlador | Coluna de Destilação |
|---|--------------------------------|-------------|----------------------|
| 1 | Temperatura do estágio 40      | TIC-100     | Deetanizadora        |
| 2 | Volume de Condensado           | LIC-101     | Deetanizadora        |
| 3 | Temperatura do estágio 16      | TIC-102-2   | Debutanizadora       |
| 4 | Vazão de Refluxo               | FIC-101-2   | Debutanizadora       |

A Tabela 2 esquematiza as variáveis secundárias escolhidas para compor a entrada do módulo neural. A coluna Controlador, informa a tag do controlador PID associado a variável secundária da planta simulada.

### 3.1.3 Estrutura do modelo

Mapear as relações dinâmicas entre variáveis primárias e variáveis secundárias do processo em estudo, é o objetivo da rede neural do sistema de inferência proposto. A arquitetura a ser utilizada é a perceptron de múltiplas camadas. A estrutura de modelagem utilizada é a NNARX para o treinamento da rede, ou seja, durante o treinamento, será realimentada pelos valores de saída da planta. Nas validações e testes da rede neural do sistema, será realimentada por suas próprias estimativas dos valores das frações molares do C2 e C5.

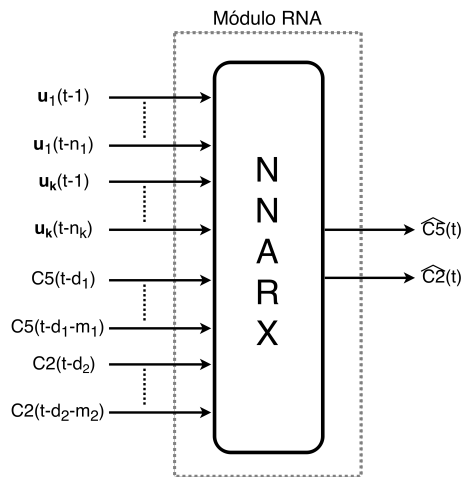


Figura 3: Diagrama esquemático da estrutura de modelagem NNARX proposta.

A Figura 3, ilustra o modelo RNA adotado para treinamento. O vetor  $\mathbf{u}$  de 4 dimensões, representa cada uma das variáveis secundárias utilizadas nesse trabalho. C2 e C5 representam o vetor de regressão constituído pelas variáveis primárias. São valores coletados do próprio sistema simulado. Já  $\widehat{C2}$  e  $\widehat{C5}$  são as variáveis primárias inferidas no treinamento da rede. O termo  $\mathbf{n}$  junto a um índice de acordo com a respectiva variável secundária representa a ordem do vetor de regressão. Da mesma forma, o termo  $\mathbf{m}$  é a ordem de regressão do vetor constituídos pelas variáveis primárias que alimentam o modelo. O termo  $\mathbf{d}$  é o atraso de cada uma das realimentações. Lembrando que o esquemático da estrutura NNARX é utilizado para treinamento do modelo. Para testes do modelo, no lugar dos próprios valores de variáveis primárias coletadas do sistema para realimentação, utiliza-se os valores que estão sendo inferidos pela rede. Assim, analisa-se o potencial do sistema operar sem a necessidade de da análise da composição do produto em um intervalo curto de tempo.

### 3.2 Módulo de Correção

Existe a possibilidade das estimativas obtidas pelo sistema de inferência divergirem acentuadamente dos valores esperados, isto é, os da simulação. Isso deve-se ao acúmulo do erro propagado a cada nova estimativa, já que variáveis primárias estimadas são utilizadas como entrada, realimentando o sistema com seus valores passados. Devido a isso, é proposto também um módulo de correção do sistema.

O módulo de correção utiliza as medições das variáveis primária realizadas pelos cromatógrafos de linha do processo, quando disponíveis ou em determinado intervalo de tempo. A ideia é corrigir o sensor virtual realizando um ajuste nas saídas da RNA. O esquemático do módulo é mostrado abaixo:

A princípio, a inferência gerada pelo módulo RNA é direcionada à saída, pois a mesma está conectada ao ponto A do esquema. No momento em que as medidas reais das variáveis primárias estiverem disponíveis pelos cromatógrafos de linha, a realimentação do sistema será feita a partir da leitura dos cromatógrafos. Para isso, a saída será conectada na posição B. Assim, corrige-se o erro acumulado com a realimentação das inferências, colocando a leitura real na saída e na realimentação. Após um instante, o sistema volta a ser conectado ao ponto A, sendo realimentado por suas próprias estimativas, até que as medidas reais estejam disponíveis novamente.

No presente trabalho, o período de medição dos cromatógrafos de linha presentes no processo é de 14 minutos. Apenas um cromatógrafo não é hábil para fornecer informações sobre diversas linhas de produção simultaneamente. Na unidade de processamento de gás natural objeto de estudo deste trabalho, a UPGN-II GMR, são medições de grande interesse a composição do GLP, do fundo e do topo da coluna deetanizadora e da carga da planta. Como há apenas um único cromatógrafo, é necessária a realização de chaveamento entre as quatro linhas do processo seguindo uma estratégia de ordenação pré-determinada. Se a sequência for serial, o tempo para obtenção da composição do GLP é de 56 min, uma vez que além de levar 14 min quando o cromatógrafo o analisa, ainda espera-se 42 min devido a análise em outras linhas.

Após a descrição dos módulos individualmente, pode-se apresentá-lo por completo. O diagrama esquemático simplificado ilustrado pela Figura 5:

Na Figura 5,  $\mathbf{VPs}$ ,  $\mathbf{C}$ ,  $\sim\mathbf{C}$ ,  $\overline{\mathbf{C}}$  são vetores e correspondem, respectivamente, às variáveis de processo dos controladores da UPGN, às medições de frações molares fornecidas pelos cromatógrafos, às frações molares estimadas pela rede neural

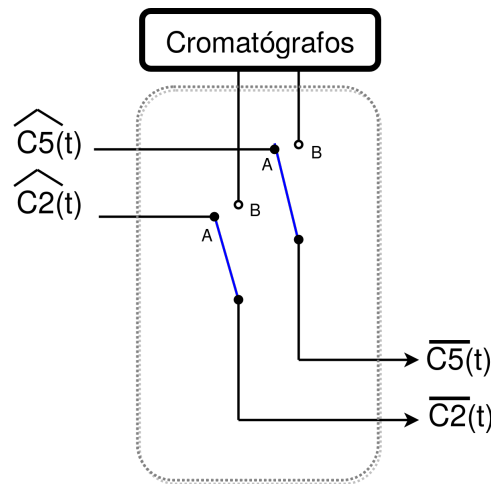


Figura 4: Diagrama esquemático do módulo de correção.

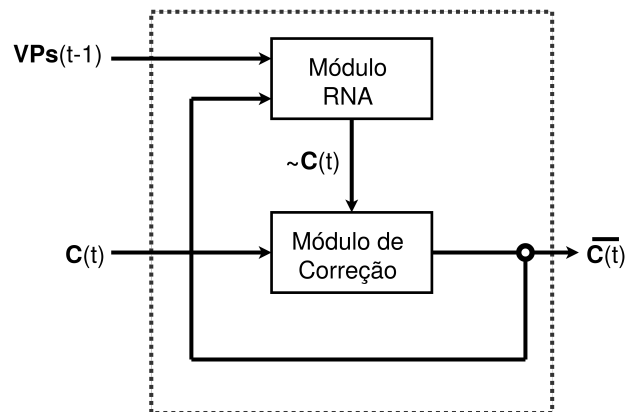


Figura 5: Esquemático do sensor virtual proposto.

e, por fim, as variáveis primárias corrigidas ou não pelo módulo de correção.

### 3.3 Implementação

Descreve-se de forma resumida a metodologia que foi utilizada no desenvolvimento da proposta deste trabalho. Apresenta-se as etapas realizadas para implementação do sistema de inferência e, consequentemente, para obter resultados.

1. **Coleta de Dados:** Na etapa de coleta de dados, obtêm-se todos os dados necessários para treinamento e validação da rede neural do modelo proposto. Para que isso seja possível, aplica-se sinais PRS aos *setpoints* dos controladores do sistema simulado, listados na Tabela 2. Coleta-se três conjuntos de dados que abragem adequadamente a faixa de operação da planta.
2. **Estrutura de inferência RNA:** Uma vez que os dados foram coletados, realiza-se a análise para diversas configurações de estruturas neurais. Treina-se a rede neural de acordo com a configuração em avaliação: estruturas de inferência de segunda ordem com diversos valores de neurônios na camada oculta. Dessa forma, pretende-se observar como a arquitetura da rede pode influenciar na sua capacidade de inferir os resultados. Depois de realizada a validação das estruturas, a melhor é escolhida para compor o sistema de inferência. A comparação entre as estruturas pode ser realizada analisando-se a complexidade da rede neural e a capacidade de inferência ótima da fração molar dos contaminantes do GLP. Neste trabalho, essa análise pode ser feita pelo erro médio quadrático da inferência.
3. **Aplicação no sistema:** Inicialmente, verifica-se o comportamento do sistema sem a presença do módulo de correção. A ideia é verificar até que momento as medições fornecidas pelo modelo são válidas para o propósito do trabalho. Após isso, utiliza-se o módulo de correção, testando-se diferentes períodos de atuação do módulo.

## 4. RESULTADOS

### 4.1 Treinamento e Validação das RNAs de inferência

Para treinamento das redes neurais que compõe o sistema de inferência, utilizou-se o *nnTool*, o *toolbox* de redes neurais artificiais do *software* MATLAB. As redes foram treinadas utilizando-se o algoritmo de treinamento: gradiente conjugado escalonado, e aplicando-se a validação cruzada e parada antecipada. De acordo com Demuth (et al., 2009), o algoritmo de gradiente conjugado escalonado funciona de forma eficiente em conjunto com a validação cruzada, justificando-se assim, sua escolha. Vale ressaltar também que o pré-processamento manual do conjunto de dados, isto é, a normalização, não foi necessária pois o próprio *toolbox* já normaliza os dados antes do treinamento e faz a operação inversa para exibir os dados de saída.

Fixou-se o número de camadas ocultas em um e realizou-se vários treinamentos para diversos valores de neurônios na camada oculta. Logo depois de cada treinamento, submeteu-se as redes neurais obtidas à 2 conjuntos de dados de validação, objetivando-se analisar sua capacidade de generalização. Assim, apresenta-se a Tabela 3 com as estatísticas de treinamento e validação.  $T_t$  representa o total de treinamentos, assim como  $T_v$  o total de validações.

Tabela 3: Estatísticas de treinamentos e validações.

| Ordem | Redes | Treinamentos | Validações | $T_t$ | $T_v$ |
|-------|-------|--------------|------------|-------|-------|
| 2     | 4     | 8            | 2          | 32    | 64    |

### 4.2 Análise das Estruturas Neurais

Apresenta-se os resultados obtidos de cada uma das estruturas de acordo com suas respectivas ordens e número de neurônios. O objetivo é analisar que estrutura obteve melhor desempenho de acordo com o seu EMQ. Assim, torna-se possível selecionar a melhor estrutura e indicá-la a compor o sistema de inferência.

Tabela 4: Estatísticas de validações da inferência de pentano.

| Ordem | $N_{NE}$ | EMQ pentano | EMQ pentano (%) |
|-------|----------|-------------|-----------------|
| 2     | 8        | 1.3991e-06  | 4.2612          |
| 2     | 16       | 2.9241e-06  | 6.7503          |
| 2     | 20       | 5.1732e-06  | 7.4512          |
| 2     | 25       | 9.0721e-06  | 8.01782         |

Tabela 5: Estatísticas de validações da inferência de etano.

| Ordem | $N_{NE}$ | EMQ etano  | EMQ etano (%) |
|-------|----------|------------|---------------|
| 2     | 8        | 1.0891e-06 | 0.4301        |
| 2     | 16       | 5.035e-05  | 3.4511        |
| 2     | 20       | 3.9915e-05 | 2.7054        |
| 2     | 25       | 2.5563e-05 | 2.4073        |

Nas Tabelas 4 e 5, exibe-se as médias dos EMQs obtidos a partir das duas validações as quais submeteu-se as melhores redes onde  $N_{NE}$  representa o número de neurônios da camada oculta. A partir da análise dessa tabela, a rede composto por 8 neurônios na camada oculta mostrou os melhores resultados, isto é, os menores EMQs tanto na inferência do pentano como na do etano. Portanto, escolhe-se a rede de 8 neurônios para formar o sistema de inferência.

Aplicou-se a rede neural escolhida a planta simulada, em tempo real, e inferiu-se os valores de pentano e etano durante um intervalo de 1500 amostras.

Na Figura 6, gráfico (a) representa a variação dos *setpoints* da temperatura do estágio 16, o gráfico (b) representa a variação nos *setpoints* da vazão de refluxo, o gráfico (c) representa a variação no *setpoint* na temperatura do estágio 40 e, por fim, gráfico (d) que represente a variação no *setpoint* do do percentual de volume do condensado. Mostra-se os resultados obtidos, sem aplicação do módulo de correção:

Na Figura 7, o gráfico (a) representa as frações molares de pentano, sendo o gráfico azul o valor gerado pelo sistema de inferência sem o módulo de correção e o vermelho os obtidos da planta real. Observa-se que nesse caso, a inferência divergiu dos valores esperados. O gráfico (b) representa as frações molares de etano, como mesmo esquema de cores da representação do pentano.

Agora aplica-se o módulo de correção para intervalos de 60 minutos, 100 minutos, 200 minutos e 300 minutos. Vale lembrar que os cromatógrafos de linha só poderão fornecer valores das frações molares a cada 56 minutos no processo real. Assim, tem-se as Figuras 8 e 9 mostrando os resultados, com a utilização do módulo, da inferência do etano e pentano, respectivamente.

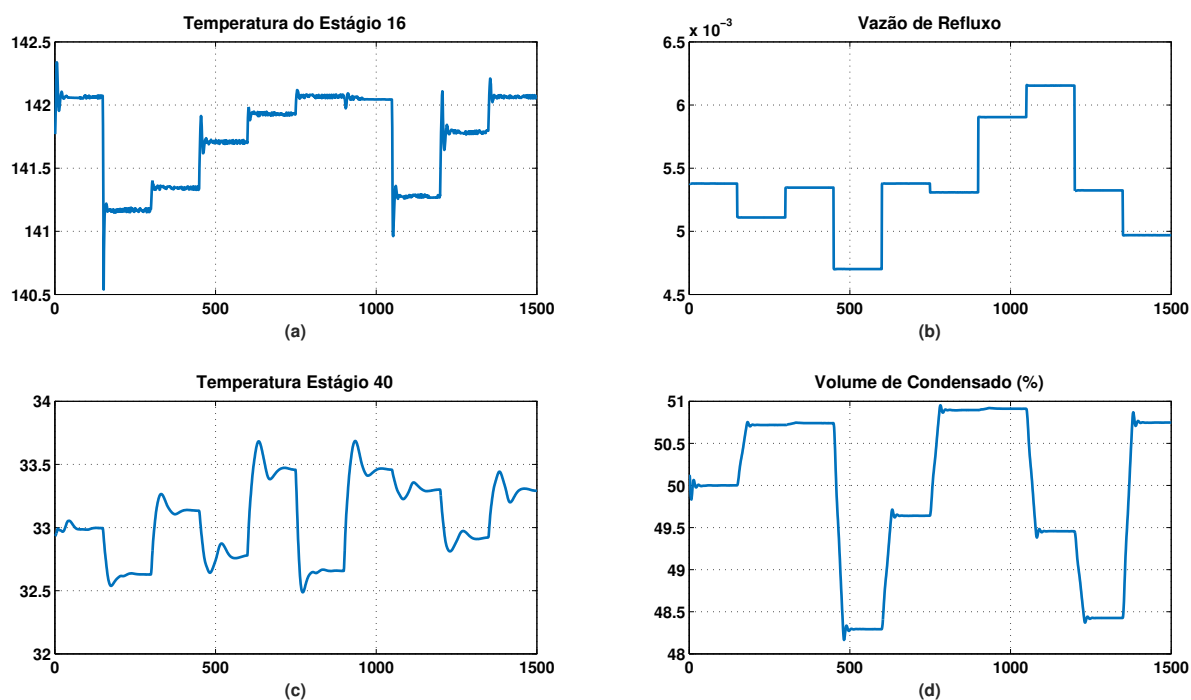


Figura 6: Esquemático dos valores das variáveis secundárias aplicados a planta simulada.

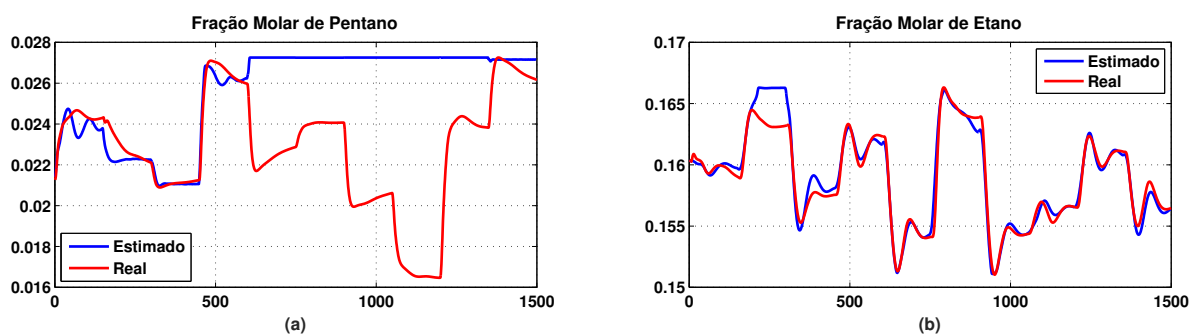


Figura 7: Resultado da inferência do pentano (a) e etano (b) sem módulo de correção.

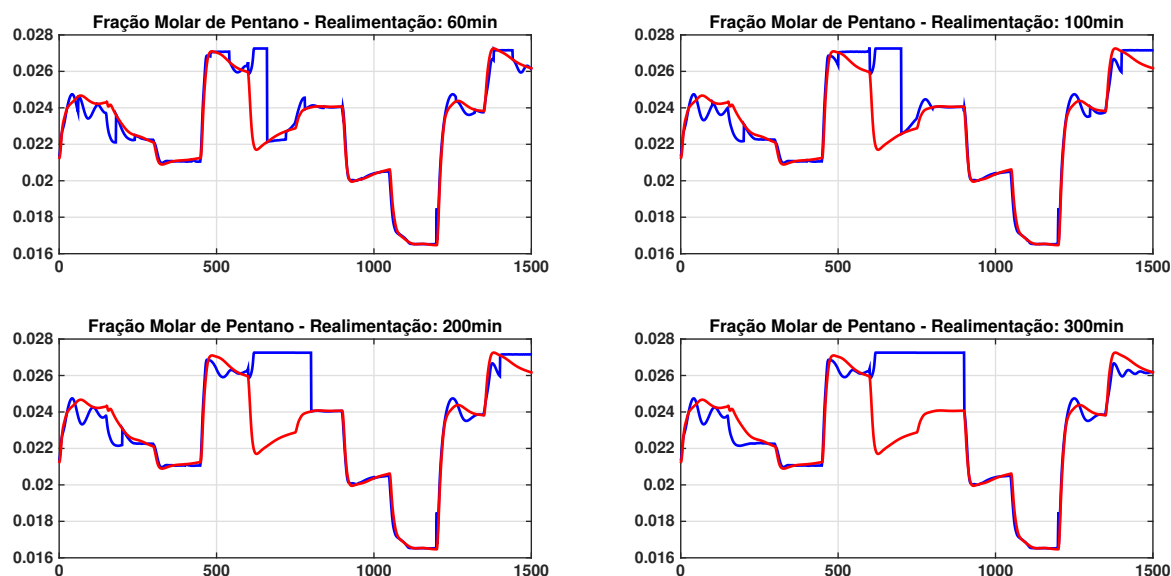


Figura 9: Inferência do pentano com o uso do módulo de correção.

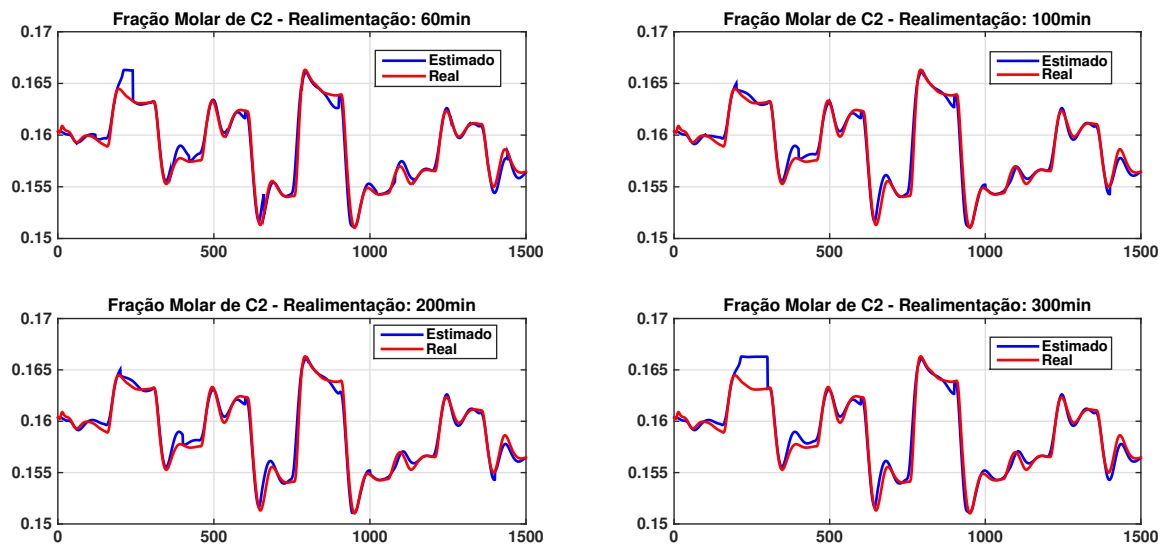


Figura 8: Inferência do etano com o uso do módulo de correção.

Ao utilizar-se o tempo mínimo de realimentação, isto é, 60 minutos, para corrigir as estimativas dos sensores, obteve-se as melhores inferências. Até mesmo com o tempo máximo, ou seja, 300 minutos para realimentar os sensores, pode-se considerar os resultados satisfatórios, já que as estimativas não divergiram dos valores esperados. Como o cromatógrafo de linha da planta pode fornecer valores das frações molares dos componentes em aproximadamente 56 minutos, se realimentarmos o sistema a cada 300min, poderia se reduzir o a quantidade de leituras dos cromatógrafos em até 6 vezes.

## 5. CONCLUSÃO

Neste trabalho, focou-se em implementar um sistema de inferência baseado em redes neurais artificiais que fosse capaz de inferir as frações molares dos contaminantes do GLP, isto é, pentano e etano.

De acordo com os resultados expostos, pode-se concluir que obter um sistema de inferência que seja eficiente é totalmente possível. Pode-se estimar frações molares do pentano e etano com boa aproximação. Na melhor configuração de rede indicada, o EMQ relativo a fração molar do pentano foi de  $1.3991e^{-06}$  sendo o seu EMQ percentual de 4.26% e do etano foi de  $1.0891e^{-06}$  sendo seu EMQ percentual de 0.4301%. Ainda considerando o pior modelo de rede, o EMQ foi de  $9.0721e^{-06}$  como EMQ percentual de 8.0178%. Já para o etano, no pior caso obtido, tem-se EMQ  $2.5563e^{-06}$  e EMQ percentual de 2.4073%. Aplicando-se a melhor estrutura obtida para inferência do C5 e C2 na planta simulada, em tempo real, mostrou-se resultados sem a aplicação do módulo de correção e com ele pra intervalos de tempo variados. Com realimentações realizadas a cada 300minutos, a inferência não divergiu dos valores reais. O melhor caso, com realimentações a cada 60 minutos, é realizável já que o cromatógrafo de linha disponibilizada os valores das frações molares a cada 56 minutos. Assim, atestou-se a aplicabilidade do sistema de inferência no processamento do gás natural, mais especificamente, para monitoramento da qualidade do GLP.

## 6. REFERÊNCIAS

- Abdalla, O.A., Zakaria, M.N., Sulaiman, S. and Ahmad, W.F.W., 2008. "Predictions of isolate and normal pentene of debutanizer catalytic reforming unit by artificial neural network". IEEE.
- Aguirre, L.A., 2004. *Introdução à identificação de sistemas: técnicas lineares e não lineares aplicadas a sistemas reais*. Editora UFMG, Belo Horizonte/MG, Brasil, 2nd edition.
- Bolf, N., Ivandic, M. and Galinec, G., 2008. "Soft sensors for crude distillation unit product properties estimation and control". In *16th Mediterranean Conference on Control and Automation*. IEEE, pp. 1804–1810.
- da Silva Linhares, L.L., 2010. *Sistema Híbrido de Inferência Baseado em Análise de Componentes Principais e Redes Neurais Artificiais Aplicado a Plantas de Processamento de Gás Natural*. Master's thesis, Universidade Federal do Rio Grande do Norte, PPgEEC, Natal/RN - Brasil.
- Demuth, H., Beale, M. and Hagan, M., 2009. *Neural Network Toolbox 6 User's Guide*. Mathworks, Natick/MA, EUA.
- Fortuna, L., Graziani, S. and Xibilia, M.G., 2005. "Soft sensors for product quality monitoring in debutanizer distillation columns". In *Control Engineering Practice*. pp. 499–508.
- Haykin, S., 2001. *Redes Neurais: Princípios e Prática*. Bookman, Porto Alegre/RS - Brasil, 2nd edition.
- Hongmei, R., Xuemin, T. and Lianfang, C., 2015. "A new soft sensor method for dynamic processes based on dynamic orthogonal forward regression". In *27th Chinese Control and Decision Conference*. IEEE, pp. 536–541.

- Linhares, L.L.S., Júnior, J.M.A. and Araújo, F.M.U., eds., 2007a. *Anais do VIII Simpósio Brasileiro de Automação Inteligente (SBAI)*, Vol. 1. Departamento de Computação e Automação, UFRN, Florianópolis.
- Linhares, L.L.S., Junior, J.M.A. and de Araújo, F.M.U., 2007b. "Redes neurais artificiais para identificação da fração molar de pentano na composição do glp". In *Anais do VIII Simpósio Brasileiro de Automação Inteligente (SBAI)*. Florianópolis/SC - Brasil.
- Norgaard, M., Ravn, O., Poulsen, N.K. and Hansen, L.K., 2001. *Neural Network for modelling and control of dynamic systems*. Springer-Verlag London Limited, Londres - Inglaterra.
- THOMAS, J., 2001. *Fundamentos de Engenharia de Petróleo*. Interciência, Rio de Janeiro/RJ - Brasil.
- Warne, K., Prasad, G., Rezvani, S. and Maguire, L., 2004a. "Development of a hybrid pca-anfis measurement system for monitoring product quality in the coating industry". In *IEEE International Conference on Systems, Man and Cybernetics 4*. pp. 3519–3524.
- Zamproga, E., Barolo, M. and Seborg, D., 2005. "Optimal selection of soft sensor inputs for batch distillation columns using principal component analysis". In *Journal of Process Control*. pp. 39–52.

## Inference System of The Molar Fraction of The Main Contaminants of Liquefied Petroleum Gas Based on Artificial Intelligence Techniques

Jean Mário Moreira de Lima, jean@dca.ufrn.br<sup>1</sup>

Fábio Meneghetti Ugulino de Araújo, meneghet@dca.ufrn.br<sup>1</sup>

<sup>1</sup>UFRN/CT, DCA - Laboratório de Controle de Processos, Lagoa Nova, Natal/RN - Brasil, CEP 59.078-970

**Resumo:** In NGPUs (Natural Gas Processment Units), one of the most profitable products is the LPG (Liquefied Petroleum Gas) which is composed mostly for propane (C3) and butane (C4). Also, pentane (C5) and etane (C2) could be present in the LPG as contaminants. The measurement of the molar fraction of the LPG's components has been done through gas chromatographs. However, chromatography is a slow process, and this makes impossible a real time quality monitoring and control of the product. In this work, it is proposed an inference system, using artificial neural network, which the main objective is inferring the molar fractions of C5 and C2. In this way, it would be possible estimate the molar fractions of those two LPG's contaminants minute by minute, which is considerable faster than the tradicional way by using cromatography. The obtained results are promising. They show that the inference system is capable to infer molar fractions of C5 and C2, then it helps to improve the LPG's quality monitoring and, consequently, profitability.

**Palavras-chave:** NGPU, LPG, Inference System, Artificial Intelligence, Artificial Neural Network