

3

Controle Preditivo Aplicado a Poços Inteligentes

3.1.

Introdução

Neste capítulo são apresentados os conceitos gerais de controle preditivo baseado no modelo (*Model Predictive control* MPC) com uma abordagem de dimensão finita no sentido de processos de Markov. Depois foram abordados os conceitos de reservatórios petrolíferos e poços inteligentes.

3.2.

Controle Preditivo: Abordagem Markoviano

3.2.1.

Controle Automático

O Controle Automático vem desempenhando um papel vital no avanço da engenharia e da ciência. Além de sua grande importância nos sistemas de veículos espaciais, no controle de mísseis e robótica, o controle automático se tornou uma ferramenta importante nos processos industriais e de fabricação. Com os avanços na teoria e prática, o controle automático proporciona os meios para conseguir um comportamento ótimo dos sistemas dinâmicos, melhorando a produtividade, simplificando trabalhos repetitivos, dentre outras atividades. (Ogata, 2003).

O objetivo em um sistema de controle consiste em aplicar sinais adequados na entrada de um processo com o intuito de fazer com que o sinal de saída do processo satisfaça certas especificações e apresente um comportamento desejado. A eficiência do processo pode ser analisada em função da técnica de controle utilizada. A Figura 8 mostra o diagrama de blocos de um sistema de controle em malha fechada.

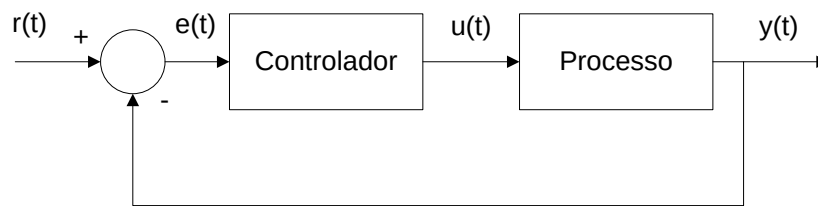


Figura 8 – Diagrama em blocos de um sistema em malha fechada.

Onde:

$r(t)$ = representa a sinal de referência para o controlador.

$e(t)$ = representa a sinal de erro entre $r(t)$ e $y(t)$.

$u(t)$ = representa a sinal de controle (lei de controle).

$y(t)$ = representa a saída do processo (variável controlada).

3.2.2. Controle Preditivo

O Controle Preditivo baseado em modelo ou *Model Predictive Control* (MPC) se refere a uma classe de algoritmos computacionais de controle que se baseia em um modelo explícito de processo para obter previsões da resposta futura da planta. A cada intervalo de controle o algoritmo MPC determina a sequência de variáveis manipuladas $u(t)$, $u(t+1)$, ..., $u(t+h)$ ajustadas tais que otimizem o comportamento futuro do processo. (Badgwell e Qin, 2001) .

O sucesso da tecnologia MPC como um paradigma pode ser atribuído a três fatores importantes: (i) incorporação de um modelo explícito da planta no cálculo de controle (Rawlings, 2000); (ii) o algoritmo MPC considera o comportamento da planta sobre um horizonte futuro no tempo, isto significa que os efeitos da retroalimentação dos distúrbios podem ser antecipados e removidos e (iii) o controlador MPC considera a entrada do processo, estados e restrições na saída, diretamente no cálculo do controle. Isto significa que as violações das restrições são muito menos prováveis, resultando em um controle mais rigoroso nas restrições ótimas de estado estacionário para o processo. Contudo, o MPC possui bom desempenho quando o processo é instável ou possui retardos.

As técnicas convencionais do MPC no horizonte finito usam um horizonte de controle, um horizonte de predição, e uma parte de *desempenho* (minimização da função custo). O horizonte de controle é o horizonte sobre o qual o agente

encontra ações. O horizonte de previsão é o horizonte sobre o qual o agente prediz o comportamento autônomo do sistema. O desempenho especifica o que o agente irá obter do estado, no final do horizonte de previsão. Normalmente, ambos horizontes são muito mais curtos do que o infinito, sendo que o horizonte de predição maior do que o horizonte de controle. O processo do agente MPC é apresentado na Figura 9

Ao considerar o horizonte de predição maior que o horizonte de controle, o agente pode analisar onde é que o sistema acaba depois que o agente executou suas ações ao longo do horizonte de controle. O princípio de horizonte finito no MPC é descrito na eq. (3.17).

$$\max_{a_{t_0}, \dots, a_{t_0+N_u}} \left[E \left\{ \sum_{k=t_0}^{t_0+N_u} r_k \right\} + E \left\{ \sum_{k=t_0+N_c+1}^{t_0+N_2} r_k \right\} + V(x_{t_0+N_2+1}) \right] \quad (3.17)$$

Onde V é a função desempenho, que indica a soma prevista de desempenho futuro que se poderá ter em um determinado estado. Em geral, essa função não é conhecida com antecedência. Pode-se supor que a função valor é zero, uma aproximação feita com uma função de Lyapunov (Jadbabaie, Yu e Hauser, 1999), ou ser aprendido com a experiência (Negenborn, Schutter, *et al.*, 2004).

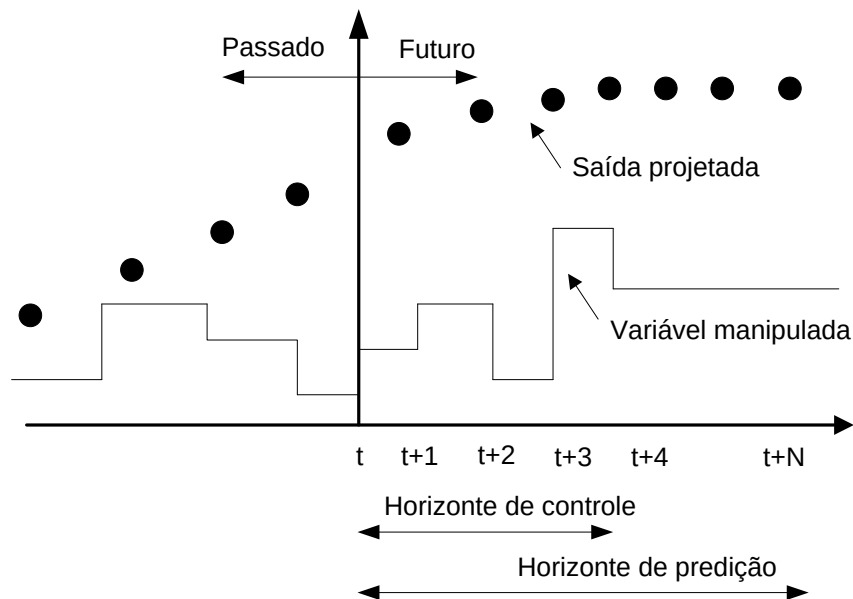


Figura 9 – Estratégia do controle preditivo.

Muitos processos são não-lineares em diferentes graus de severidade. Embora existam muitas situações em que o processo possa ser operado na vizinhança do estado estacionário e, portanto, uma representação linear poderia ser adequada, existem algumas situações importantes onde isto não acontece (F. Camacho e Bordons, 1999). Por outro lado existem processos nos quais as não-linearidades são muito severas e cruciais na estabilidade da malha fechada, onde a aplicação de um modelo linear é insuficiente. Nesses casos, os modelos não-lineares poderiam ser essenciais para melhorar o funcionamento ou simplesmente para operar de modo estável. Não há nada nos conceitos básicos da MPC contra o uso de um modelo não-linear, portanto, a extensão de idéias de MPC para processos não-lineares é simples, pelo menos conceitualmente. Contudo, é um problema não trivial com muitas questões a serem esclarecidas:

- A disponibilidade de modelos não-lineares, devido à falta de técnicas de identificação para processos não-lineares.
- A complexidade computacional para resolver modelos preditivos de controle em processos não-lineares.
- A falta de resultados em estabilidade e robustez para o caso de sistemas não-lineares.

Alguns destes problemas são parcialmente resolvidos. Desta maneira, o MPC associado ao uso de modelos não-lineares vem despertando um grande interesse e tem sido bastante estudado (Morari e Lee, 1999). Hoje em dia o controle preditivo não-linear ou *Nonlinear Model Predictive Control* (NMPC) ocupa as novas pesquisas do MPC; eventos como *International Workshop on Assessment and Future Directions of Nonlinear Model Predictive Control*, relatam as pesquisas e aplicações na área da petroquímica, polímeros, plantas de gás, etc (Badgwell e Qin, 2001). Estes novos modelos NMPC se tornaram mais comum entre usuários que exigem maior desempenho.

3.2.3. Controle Preditivo Neural

O Controle Preditivo Neural ou Neural Generalized Predictive Control (NGPC) utiliza as redes neurais para a modelagem do processo a ser controlado (Figura 10). Ele consiste de quatro componentes: (i) a planta a ser controlada; (ii) os modelos de referência que especificam o desempenho desejado da planta; (iii) a rede neural que modela a planta e (iv) o algoritmo de minimização da função de custo que determina a entrada necessária para que se produza a saída otimizada. Portanto, o algoritmo NGPC consiste em um bloco de minimização da função custo e o bloco da rede neural (Soloway e Haley, 1996).

O funcionamento do NGPC começa com um sinal de entrada $r(t)$ o qual é apresentado ao modelo de referência. Este modelo produz uma trajetória do sinal de referência y_m a ser seguida, a qual é usada como entrada ao bloco de minimização da função custo. O algoritmo de minimização da função custo produz uma saída que é usada tanto como uma entrada na planta quanto no modelo da planta. A saída predita $y_n(t)$ do modelo da planta é comparada com a trajetória de referência cujo erro produzido é então enviado ao bloco do otimizador, que contém uma função custo. O algoritmo de minimização da função custo produz uma saída utilizada como entrada para a planta $u(t)$ quando esta é minimizada, e o processo se repete. A chave na saída de $u(t)$ só vai para a planta quando o melhor $u(t)$ for encontrado, no entanto a chave vai ate o modelo da planta (rede neural), onde o algoritmo de minimização de custo utiliza este modelo para calcular a seguinte entrada de controle $u(t + 1)$ desde as predições da resposta do modelo da planta conforme ilustrado na (Figura 10). Podemos ter restrições observadas na função custo que limitam o grau de liberdade das variáveis.

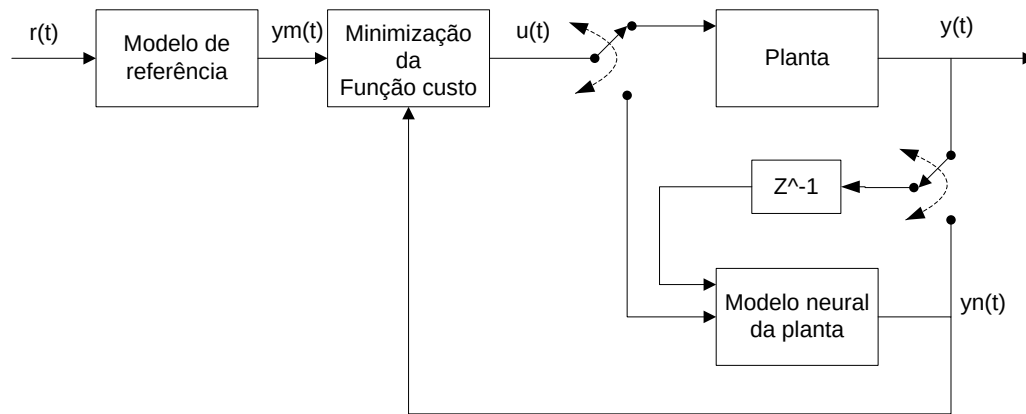


Figura 10 – Diagrama em bloco do controlador neural preditivo.

3.2.3.1. Função de Custo

O algoritmo NGPC é baseado na minimização da função custo sobre um horizonte de predição finito. Em geral, a função custo em um MPC e NGPC segue a eq.(3.18).

$$J = \sum_{j=N_1}^{N_2} [y_m(t+j) - y_n(t+j)]^2 + \sum_{j=1}^{N_u} \lambda [\Delta u(t+j)]^2 \quad (3.18)$$

Onde:

- N_1 = Horizonte mínimo de predição.
- N_2 = Horizonte máximo de predição.
- N_u = Horizonte de controle.
- $Y_n(t+j)$ = Predição do sinal de saída do processo.
- $Y_m(t+j)$ = Trajetória de referência ao longo do tempo.
- λ = Fator de ponderação da entrada de controle.
- $\Delta u(t+j)$ = Variação entre $u(t+j)$ e $u(t+j-1)$.

A função custo pode ser calculada diretamente, resultando em uma expressão de ação de controle:

$$\tilde{u} = (G^T G + \lambda I)^{-1} G^T (y_m - f) \quad (3.19)$$

Onde $\tilde{u}$ é um vetor que contém as mudanças na ação de controle, I é a matriz identidade, G é uma matriz que baseia a resposta em relação ao degrau do processo, sendo utilizado para calcular o efeito que as ações de controle do futuro terão na saída do processo e, f é um vetor que contém as previsões futuras da saída do processo até o horizonte N_2 . A função custo (eq (3.18)) permite minimizar não somente o erro quadrático médio existente entre o sinal de referência e o modelo da planta, como também a taxa da variação quadrática da entrada de controle (Soloway e Haley, 1996).

O emprego de um modelo não-linear, seja por redes neurais ou por modelos NARX (Nonlinear AutoRegressive with eXogenous inputs) e NARMAX (Nonlinear AutoRegressive Moving Average with eXogenous inputs), introduzem o problema que a função custo (eq. (3.18)) só pode ser minimizada através de uma abordagem iterativa, já que solução direta da eq. (3.19) não é possível, por ser só aplicável a modelos lineares (Lennox e Montague, 2001). Embora a convergência da solução não seja garantida, o uso de algoritmos de otimização não-lineares, como por exemplo Quase-Newton ou no caso do trabalho precursor de NGPC (Soloway e Haley, 1996), minimiza a função custo mediante o algoritmo de Newton-Raphson, ou ainda, os trabalhos que utilizam a inteligência computacional como algoritmos genéticos para encontrar a solução ótima multicritério (Laabidi, Bouani e Ksouri, 2008), são freqüentemente empregados.

3.2.4. MPC no Processo de Decisão de Markov

Nesta seção será avaliado como um *agente* MPC, que tradicionalmente é utilizado para controlar sistemas dinâmicos, pode ser descrito como um processo de decisão de Markov, ou *Markov Decision Processes* MDP. Esta pesquisa se baseia nos trabalhos de (Negenborn, Schutter, *et al.*, 2004)(Negenborn, Schutter, *et al.*, 2005).

Primeiramente, é preciso definir a maneira como o agente pode encontrar uma ação ótima. O agente, de alguma forma, tem que usar o sistema e o desempenho do modelo para encontrar uma sequência de ações que proporcionem o melhor desempenho ao longo do horizonte de controle finito. Do ponto de vista gráfico dos processos de decisão de Markov, esta se resume a encontrar o

caminho de passos N_u que apresenta o maior desempenho esperado acumulado. A soma total de desempenho acumulada dá como resultado a sequência com desempenho acumulado. Estas considerações conduzem-nos ao seguinte algoritmo simples, para o MPC, nos processos de decisão de Markov:

- Deslizar o horizonte para o passo atual de decisão, observando o estado do sistema, e definir o problema de otimização para encontrar as ações ao longo do horizonte de controle que maximize o desempenho a partir do estado observado.
- Encontrar todos os caminhos de comprimento N_u e acumular o caminho com o maior desempenho esperado acumulado.
- Aplicar a primeira ação desta sequência e ir para o passo seguinte de decisão.

3.2.4.1. Abordagem da Função Valor

Deixemos de lado a eq. (3.18), correspondente a um MPC o qual poderia ser chamado de modelo clássico, nesta seção iremos levar o agente MPC descrito num sentido de processos de Markov, a eq. (3.17) de um agente MPC descrito no horizonte finito, nos ajudará mais a descrever este novo modelo. A função valor V profere o premio acumulado esperado futuro para cada estado x e uma política π . O ótimo valor da função V^* apresenta o mais elevado premio possível acumulado esperado futuro para cada estado. O mais elevado premio possível futuro é obtido pelas seguintes ações que prescrevem uma política ótima π^* . A partir de agora considerar-se-á uma política probabilística.

$$V^*(x_{k_0}) = \max_{\pi} E \left\{ \sum_{k=k_0}^{\infty} r(x_k, \pi(x_k), x_{k+1}) \right\} \quad (3.20)$$

Na eq.(3.20) r é o desempenho recebido para a transição na decisão no passo k e x são os estados. A ótima função valor V^* é então obtida através da resolução de cada x_{k_0} . Assumimos que o valor da função ótima é conhecido. Do ponto de vista gráfico de um MDP, isso significa que se pode indicar cada nó com o desempenho futuro esperado. Neste caso, quando o sistema está no estado x , o

agente tem que considerar apenas as ações $a \in A_x$ possíveis no estado x e encontrar a ação que resulte a maior soma de recompensa obtida diretamente, mais a recompensa acumulada esperada futura do estado resultante após a ação ter sido executada, conforme descreve a eq. (3.21). Esta soma é chamada de Q , valor para o par (x,a) e é usada pelo agente para encontrar a ação que resulta no maior valor de Q .

$$a_k = \arg \max_{a \in A_{x_k}} \left[\sum_{x'} P(x'|x_k, a) (r(x_k, a, x') + V^*(x')) \right] \quad (3.21)$$

Assim, quando a função valor ótimo é conhecida, ao invés de considerar passos N_c , o agente tem de considerar apenas um passo do procedimento de otimização em cada processo de decisão, ou seja, o horizonte de controle torna-se $N_c = 1$. Além disso, uma vez que a função valor é ótima ao longo do horizonte infinito, as ações escolhidas também serão ótimas ao longo do horizonte infinito.

3.2.4.2. Programação dinâmica

O problema é que, em geral, nem políticas ótimas nem funções valor ótimas para o controle de sistemas dinâmicos são conhecidos antecipadamente. Assim, a questão é: como a função valor pode ser calculada? Uma vez que se está considerando o caso de horizonte infinito, o agente não pode simplesmente calcular a função valor, já que não pode explicitamente somar o desempenho ao longo de um horizonte infinito. Então, em vez disso, é necessário aproximar a função valor de alguma maneira. Uma forma de aproximar a função valor é pelo uso de um *fator de desconto*. Este fator de desconto faz com que a soma infinita de desempenho encontre um ponto de convergência. Podemos empregar métodos de *programação dinâmica* para encontrar a função de valor neste caso.

Métodos de programação dinâmica (Bellman, 1957) (Bertsekas e Tsitsiklis, 1996) aproximam o valor de uma função quando o fator de desconto é usado. Dada uma política, calcular com métodos de programação dinâmica a função valor como a eq.(3.22)

$$V(x_{k_0}) = E \left\{ \sum_{k=k_0}^{\infty} \gamma^{k-k_0} r_k \right\} \quad (3.22)$$

Onde $\gamma \in (0,1)$ é o fator de desconto. Entanto γ seja mais perto de 1 as expectativas de desempenho a longo prazo são tomadas com maior consideração (Negenborn, Schutter, *et al.*, 2004). O fator de desconto faz com que o desempenho recebido anteriormente seja mais importante do que o desempenho obtido no futuro. Podemos reescrever a função de valor como:

$$\begin{aligned} V(x_{k_0}) &= E \left\{ \sum_{k=k_0}^{\infty} \gamma^{k-k_0} r_k \right\} \\ &= E \{ r_{k_0} + \gamma r_{k_0+1} + \gamma^2 r_{k_0+2} + \dots \} \\ &= E \{ r_{k_0} + \gamma V(x_{k_0+1}) \} \\ &= \sum_{a \in A_{x_{k_0}}} P_{\Pi}(x_{k_0}, a) \left[r(x_{k_0}, a, x') + \gamma \sum_{x'} P(x'|x_{k_0}, a) V(x') \right] \end{aligned} \quad (3.23)$$

Onde $P_{\Pi}(x, a)$ é a probabilidade de que a política atribui a escolher uma ação a no estado x . A eq. (3.23) é chamada equação de Bellman a forma da função valor. Equações de Bellman relacionam as funções valor recursivamente a si mesmo. Neste caso a equação de Bellman para todos os estados de um sistema de equações escolhe a única solução que é a função valor ótima. O algoritmo proposto por (Negenborn, Schutter, *et al.*, 2004) para definir um algoritmo MPC para MDP é o seguinte:

1. Aplicar o princípio de horizonte móvel, atualizando a estimativa do estado com uma medição do estado.
2. Calcule o valor da função dado o modelo do sistema mais recente.
3. Formular o problema de otimização em um horizonte de controle $N_u = 1$ e encontrar a ação que traz o estado do sistema para o estado com maior valor. Resolver o problema de otimização.
4. Implementar a ação encontrada e passar para a próxima etapa de decisão.

A principal desvantagem está em calcular a função valor ótima para cada passo de decisão, dado o alto custo computacional exigido. Calcular a função de valor ótimo *offline* antes que o agente comece a controlar o sistema, tem a desvantagem do que o sistema não poder variar ao longo do tempo. Embora o horizonte deslizante forneça alguma robustez, as mudanças estruturais nos parâmetros do modelo de sistema não são antecipados.

Ao invés da função $V(x_{k_0})$ recalcular o valor em cada passo de decisão, pode-se atualizar a função valor utilizando a experiência da interação entre o agente e o sistema de verdade. Se propõe em (Negenborn, Schutter, *et al.*, 2004) combinar MPC para processos de decisão de Markov com aprendizagem da função valor on-line utilizando aprendizado por reforço. Dessa forma, as mudanças no sistema são antecipadas on-line, enquanto não se esteja computando a função valor em cada passo de decisão.

3.2.5. MPC com Aprendizado por Reforço

Nesta seção é apresentada a combinação do MPC para processos de decisão de Markov, com aprendizado da função valor *on-line* mediante RL (Sutton e Barto, 1998). Nesse tipo de aprendizado, há um sistema estocástico no qual não se conhece as probabilidades de transição e o desempenho das ações obtidas em estados individuais. O agente deve encontrar uma política que dá o máximo de desempenho esperado acumulado futuro, esse agente tem de interagir com o sistema, que contém implicitamente o modelo do processo para determinar o valor ótimo da função.

Procurando seguir a notação já introduzida, pode-se denotar por r_k o desempenho obtido pelo agente MPC na função custo na eq.(3.18) no passo k , por $a_0, \dots, a_\infty$ as ações de controle a serem determinadas pelo mesmo e por E o operador de esperança. Sem negligenciar a natureza estocástica do sistema em consideração, pode-se resumir formalmente a tarefa do agente MPC como sendo a de solucionar o seguinte problema de otimização:

$$\max_{a_0, \dots, a_\infty} E \left\{ \sum_{k=0}^{\infty} r_k \right\} \quad (3.24)$$

Neste trabalho foi desenvolvido um algoritmo MPC com TD (*Temporal-Difference*), o método TD, conforme mencionado no capítulo 2 é um dos métodos para calcular a função valor sem um modelo da dinâmica do ambiente (Sutton e Barto, 1998). O aprendizado por diferença temporal minimiza a diferença entre o valor estimado dos passos de decisões sucessivas explicitamente, usando o valor estimado de estados sucessivos, como indica a eq. (3.25).

$$V(x_t) = V(x_t) + \alpha x_t (r_t + \gamma V(x_{t+1}) - V(x_t)) \quad (3.25)$$

O controle ótimo no MPC é encontrado através de um processo de otimização. Quando um modelo linear é utilizado para prever a saída, o controle ótimo pode ser achado analiticamente mediante minimização quadrática da eq.(3.18), mas caso aja um modelo não-linear é necessário a aplicação de um método de otimização mais sofisticado. Possíveis soluções envolvem programação dinâmica (Kirk, 1970)(Bellman, 1957)(Bertsekas, 2005), cálculo variacional, ou programação evolutiva.

Os trabalhos de (Negenborn, Schutter, *et al.*, 2004)(Negenborn, Schutter, *et al.*, 2005) formulam o MPC em termos de Processos de Decisão de Markov e definem a função custo pelo método TD(λ). Levando em conta as considerações teóricas de (Negenborn, Schutter, *et al.*, 2005) é apresentado um modelo que aborda a combinação do MPC-RL. Trabalhos similares se encontraram em (Gaweda, Muezzinoglu, *et al.*, 2006), mas utilizam o método SARSA (Sutton e Barto, 1998) como mecanismo de otimização para a determinação de doses de droga em um modelo de controle preditivo aplicado na administração de anemia renal.

3.3.

Poços Inteligentes no Desenvolvimento de Campos Petrolíferos

O desenvolvimento de um campo petrolífero pode ser entendido como um conjunto de ações (perfurações, sistemas de injeção, plataformas, etc) necessárias para colocar o campo em produção. A forma como é feito esse desenvolvimento é uma tarefa das mais importantes na área de reservatório, dado que esta decisão

afeta o comportamento do reservatório, decisões futuras, análises econômicas e conseqüentemente, a atratividade resultante do projeto definido. Isto envolve variáveis tais como localização, números e tipos de poços, as características do reservatório e inclusive o cenário econômico.

3.3.1. Fundamentos de Simulação de Reservatórios

A simulação de reservatórios é uma das principais áreas dentro da engenharia de reservatório, onde são aplicados modelos matemáticos mediante simuladores numéricos e computacionais (Thomas, 2001) com a finalidade de analisar o comportamento de um reservatório. Um simulador de reservatório é um processo computacional através do qual o especialista em reservatórios, a partir de informações (geológicas e geofísicas) reais, obtém previsões sobre a produção de óleo, gás e água em qualquer intervalo de tempo.

Os simuladores numéricos permitem mais sofisticação nos estudos dos reservatórios, porém, é necessário dispor de dados da rocha, dos fluidos, da geologia, do histórico de produção, não só em quantidade, mas dados com boa qualidade.

3.3.2. Poços Inteligentes

Um poço inteligente é um poço não convencional com completações inteligentes. A completação de poços consiste na transformação da perfuração em uma unidade produtiva completamente equipada e com os requisitos de segurança atendidos, pronta para produzir óleo e gás, gerando receitas. As completações inteligentes podem ser definidas como completações com instrumentação instalada na tubulação de produção, a qual permite o monitoramento contínuo e o ajuste das taxas de fluxo dos fluidos e das pressões.

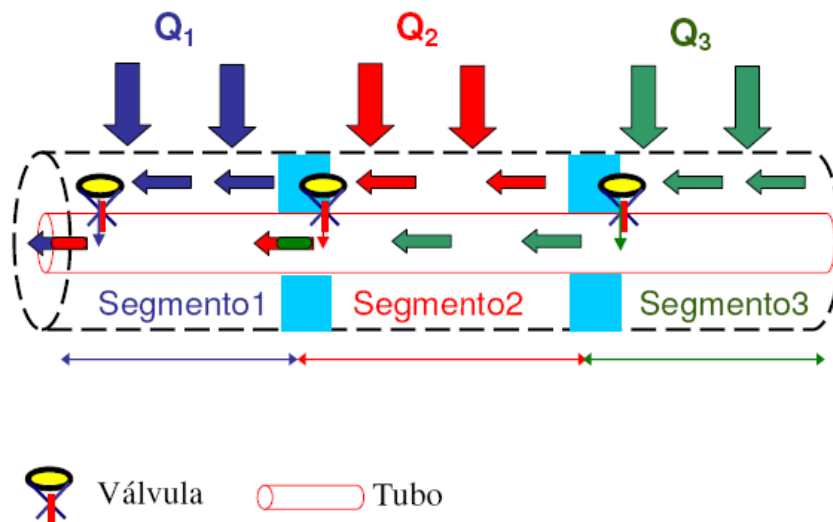


Figura 11 – Exemplo de completações de um poço horizontal Inteligente.

A Figura 11 mostra um poço horizontal inteligente dividido em três segmentos. Cada segmento contém uma válvula de controle de fluxo onde Q_1 , Q_2 , e Q_3 representam a vazão da entrada do fluxo nos segmentos 1, 2 e 3 respectivamente. Em uma visão mais global, com o desenvolvimento de poços inteligentes, isto é, de poços com instrumentação na perfuração, surge a possibilidade de alcançar um gerenciamento da produção do campo inteiro que permita, entre outras coisas, aumentar a vida útil do campo, e também dos poços. Este gerenciamento implica em realizar o controle e otimização da produção aplicando estratégias de controle de laço fechado e também otimização do processo de injeção de água (*waterflood*) com o objetivo de ter maior controle da produção de óleo, da vazão de produção de água do reservatório alongando o tempo útil dos poços (Aitokhuehi, 2004) (Faletti, 2007).

O conceito de poços inteligentes se propõe a baratear as operações de restauração mais corriqueiras, tais como o isolamento e a abertura de novos intervalos produtores, além do monitoramento em tempo real dos dados de produção, vazões, pressões e temperatura, permitindo um melhor gerenciamento do reservatório.

A completação inteligente também pode ser definida como um sistema capaz de coletar, analisar e transmitir dados para o acionamento remoto de dispositivos de controle de fluxo, com o objetivo de conseguir otimizar a produção do reservatório. Entretanto, essa tecnologia tem um custo associado

bastante alto, em geral, aliado ao fato de ser relativamente recente, sem muitos dados relativos à sua confiabilidade e melhor forma de utilização. Isto tem criado certo receio em aprovar sua implantação, especialmente por não existir ainda uma metodologia padronizada para o cálculo de seus benefícios.

3.3.3. Estratégias de Controle

O controle de produção através da operação das válvulas, existentes em completações inteligentes, pode ser feito com base em duas estratégias distintas: uma estratégia denominada como “reativa” e outra “pró-ativa” (Yeten, 2003) (Faletti, 2007).

Na estratégia de controle reativa a operação das válvulas é feita em reação a comportamentos observados na produção de fluidos dos poços. Ao se observar, por exemplo, um aumento na relação água / óleo de uma determinada região, adota-se como reação a restrição dos fluxos nessa região, privilegiando a produção em outra região que apresente menor relação água / óleo.

Por outro lado, na estratégia de controle pró-ativo, a programação da operação das válvulas é feita de forma antecipada, isto é, procura-se agir antes do efeito. Como exemplo, pode-se citar a busca, desde o início da produção de uma configuração de operação das válvulas que satisfaça alguns objetivos como: atrasar o início do corte de água nos poços, antecipar a produção, ou alcançar uma maior recuperação de óleo do campo. A estratégia de controle proativo implica em um processo de otimização com um horizonte de previsão de produção e de controle de válvulas mais longo e a necessidade de um modelo de reservatório que se ajuste a este horizonte de previsão (Faletti, 2007).

Nesta dissertação a otimização de controle de válvulas é realizada através da estratégia de controle reativa, sendo esta a estratégia que atende aos objetivos do estudo em questão.