

5 Estudos de Casos:

Neste capítulo se mostram os experimentos que foram realizados durante o desenvolvimento do modelo proposto. Estes experimentos estão organizados da forma que a Figura 34 mostra, a seguir:

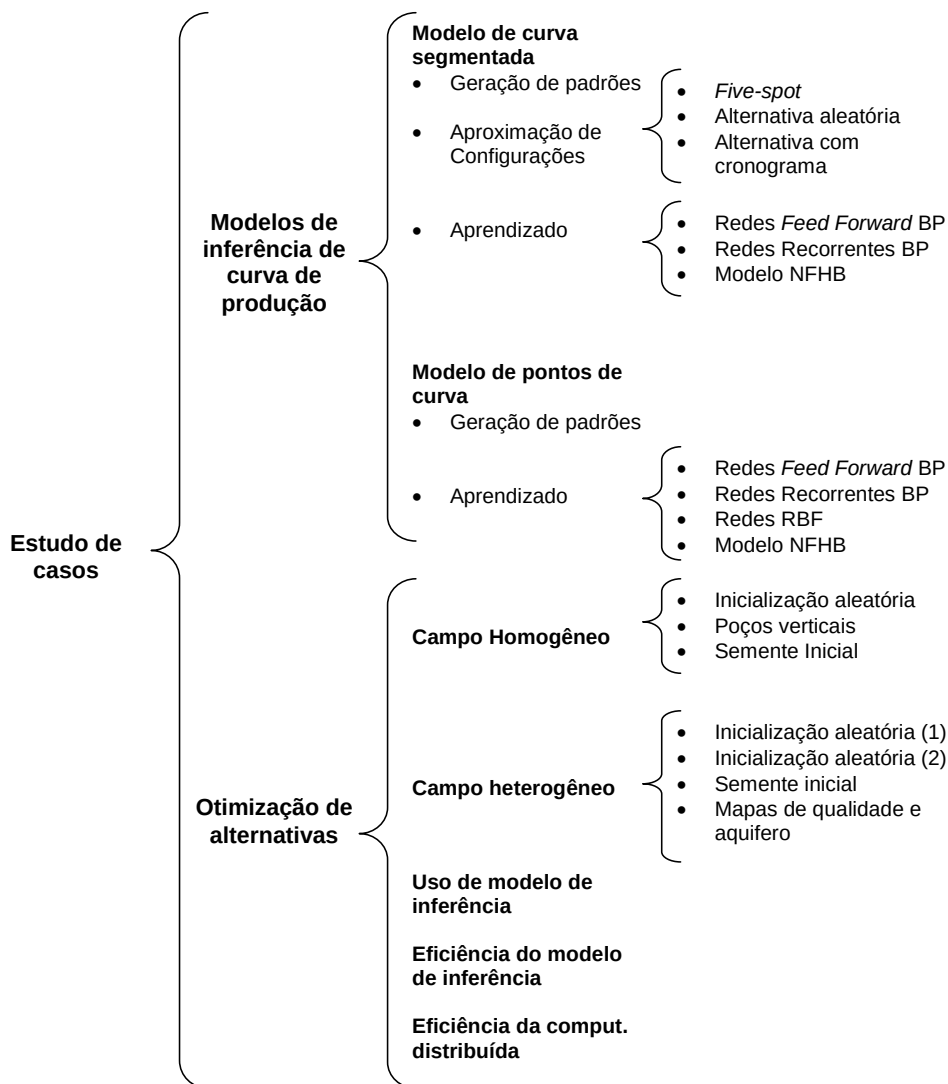


Figura 34. Taxonomia do estudo de casos deste trabalho

Como é mostrado na Figura 34, os experimentos estão divididos em experimentos dos modelos de inferência da curva de produção e experimentos da otimização de alternativas. Os experimentos dos modelos de inferência foram realizados para os dos modelos de inferência especificados neste trabalho, onde

foram criadas amostras para o aprendizado (listas de poços e seus respectivos valores de saída), e realizados os experimentos de aprendizado com os modelos aproximadores neurais e NFHB. A otimização de alternativas foi testada empregando duas configurações de reservatório uma com características homogêneas e a outra heterogênea, as quais são descritas a seguir

5.1.

Reservatório 30x30x1

Este modelo de reservatório consiste em uma malha de tipo retangular de de 30x30x1 blocos, sendo as dimensões de cada bloco 100x100x90 metros, com os seguintes valores de permeabilidade e porosidade.

permeabilidade	1000,0 (md) em todas as direções
porosidade	0,2
pressão inicial	$100 \frac{kg}{cm^2}$ em todos os blocos

Estes valores são os mesmos ao longo da reserva. Cabe ressaltar que, nesta configuração, a saturação inicial de água é 0,2 para todos os blocos, fazendo com que o potencial de óleo seja igual em todos os blocos, o que caracteriza a homogeneidade do reservatório. Na Figura 35, a seguir, mostra-se uma visão em 3D deste modelo de reservatório.

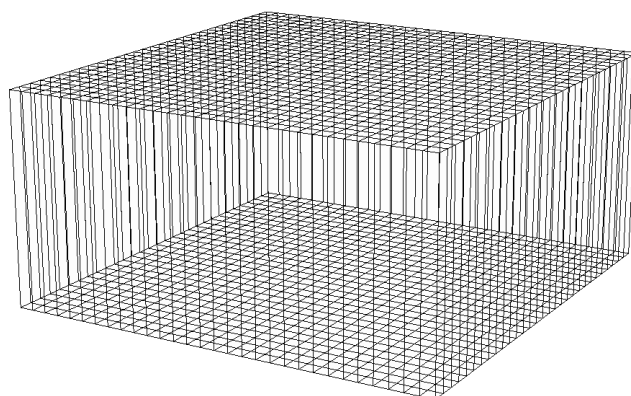


Figura 35. Modelo de reservatório 30x30x1

Com este reservatório foram realizados experimentos referidos aos modelos de aproximação da função de produção, que empregaram alternativas com poços verticais injetores e produtores e experimentos para testar o sistema de otimização.

5.2. Reservatório 33x57x3

Este modelo de reservatório consiste em uma malha de 33x57x3 blocos com dimensões dos blocos de aproximadamente 100,0x100,0x8,66 metros.

Os valores geológicos do reservatório são os seguintes.

permeabilidade	575,0 (md) nas direções i, j
	57,40 (md) na direção k
porosidade	0,23
pressão inicial	de 390 até $420 \frac{kg}{cm^2}$ dependendo inversamente da 'altura'.

A característica deste modelo é a existência de uma região com 100% de saturação de água, na parte mais funda formando um aquífero. As demais regiões possuem saturação de água no valor de 0,25. Este fato cria a condição de heterogeneidade da reserva. Na figura mostra-se a visão 3D deste campo.

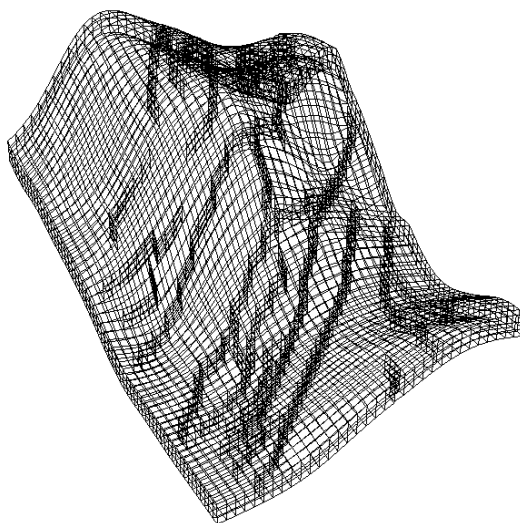


Figura 36. Modelo do reservatório 33x57x3

Neste reservatório foram realizados experimentos referidos à otimização de alternativas, aprendizado dos modelos de inferência de curvas de produção, do uso da distribuição da avaliação no ambiente de computação em paralelo e experimentos do uso de informação do mapa de qualidade.

No Apêndice B, encontra-se o conteúdo do arquivo .DAT que descreve as duas configurações de reservatório empregadas.

5.3. Aproximação da Função de produção

A realização dos testes para análise de desempenho dos modelos de aproximações da função de produção requer a seguinte seqüência:

- Teste dos modelos de curva de produção;
- Geração dos padrões de aprendizado;
- Teste do aprendizado com os modelos neurais e *neuro-fuzzy*.

5.3.1. Testes do modelo de curva segmentada

O modelo de curva segmentada foi desenvolvido neste trabalho pelo fato da curva típica de produção de óleo usada nas análises financeiras seguir este formato.

Para se obter os parâmetros do modelo que correspondem a uma curva de produção do simulador foi usada a abordagem baseada no algoritmo genético.

O algoritmo genético usado é do tipo clássico (não híbrido) com os seguintes operadores:

- cruzamento $\Rightarrow$ simples, aritmético e geométrico;
- mutação $\Rightarrow$ uniforme, não uniforme, de fronteira e gaussiana.

Como a curva segmentada possui restrições de domínio, estas restrições foram estabelecidas conforme descrito na Tabela 5 abaixo.

Tempo	Produção Óleo	Amortecimento
$t_a \in [0.1, 365]$	$q_a \in [0, 40000]$	$\rho \in [0, 1.0]$
$t_b \in [0.2, 730]$	$q_b \in [0, 40000]$	
$t_c \in [0.3, 7300]$		

Tabela 5. Restrições existentes nos parâmetros do modelo segmentado

As restrições lineares existentes também foram consideradas, como mostram as equações (25),(26),(27) da seção 4.5.2.2.

Os parâmetros de evolução do algoritmo genético empregado foram os seguintes:

Parâmetro	Valor
Tamanho da população	100 indivíduos
Nº de gerações	100 gerações
Nº de rodadas	2 rodadas
Nº de ciclos	1 ciclo

Tabela 6. Parâmetros da evolução – modelo segmentado

O tempo necessário para a execução completa de cada uma das otimizações foi inferior a 1 segundo.

Nas seções seguintes, apresentam-se experimentos da aproximação da curva segmentada para três configurações de poços: o Experimento 1, que se refere à aproximação para uma configuração 4 *five-spots*; o Experimento 2, que aproxima a curva de uma alternativa não ótima; e o Experimento 3, o qual aproxima a curva de uma alternativa 4 *five-spots* com cronograma de poços.

5.3.1.1.

Experimento 1: 4 *five spots*

Este experimento foi realizado com uma configuração de poços típica da engenharia de reservas denominada 4 *five-spots*. Esta organização consiste em colocar 4 configurações *five-spot* (um poço vertical de um tipo rodeado de 4 poços verticais do outro tipo) distribuídas no campo petrolífero, como se mostra na Figura 37 a seguir.

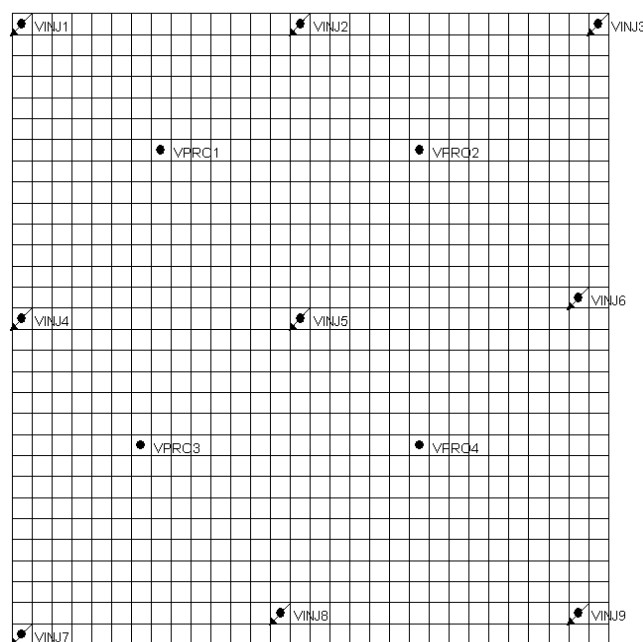


Figura 37. Reservatório com a configuração 4 *five-spots*.

Na Figura 38 e na Figura 39 mostram-se a curva de produção de óleo acumulado e a curva de produção de óleo diária, respectivamente, para a configuração de poços mostrada na Figura 37, que foram obtidas através de simulação.

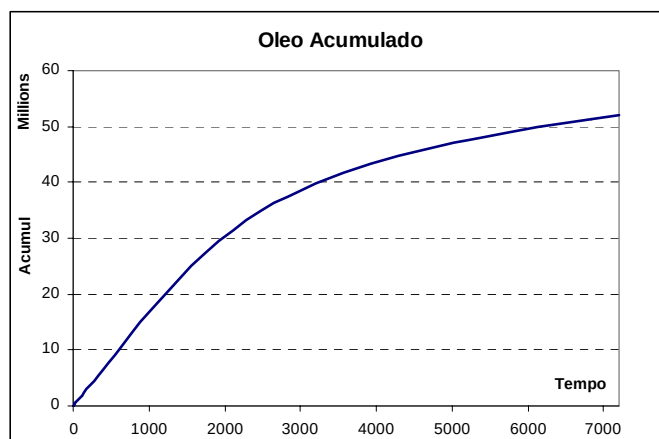


Figura 38. Curva de óleo acumulado: 4 *five-spots*

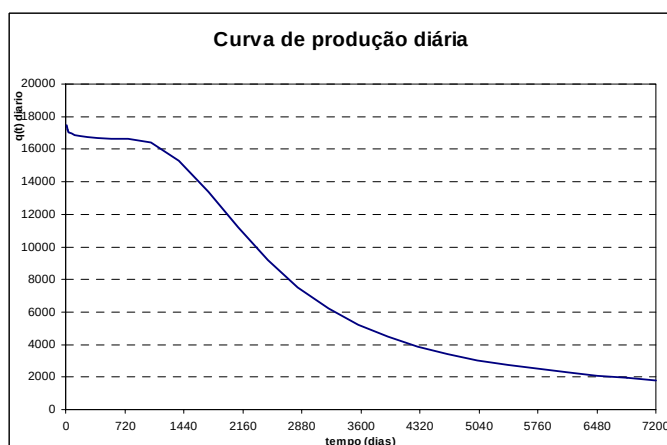


Figura 39. Curva de produção de óleo: 4 *five-spots*

O experimento foi realizado, utilizando como referência, a curva de produção de óleo diária da Figura 39. Durante a execução do algoritmo genético foram empregadas as taxas de operadores da Tabela 7, a seguir.

Taxas de operadores		Valor
Taxa de <i>crossover</i>		0.65
Taxa de mutação		0.08

Tabela 7. Taxas de operadores genéticos – modelo segmentado

Após a execução da otimização pelo AG, a melhor solução encontrada é a seguinte, mostrada na Tabela 8 a seguir.

t_a	t_b	t_c	q_a	q_b	ρ	MAPE
1,421	3,726	1097,6	16800	16800	0,169	1,66%

Tabela 8. Parâmetros obtidos para o Experimento 1: 4 *five spots*.

onde podem ser vistos os parâmetros da solução e também o erro MAPE calculado em 1,66% entre a curva de produção do modelo segmentado com estes parâmetros e a curva original gerada pelo simulador IMEX.

A partir dos valores de t_a , t_b , t_c , q_a , q_b e ρ , obtém-se a curva resultante aplicando a função segmentada da equação (23), e utilizando os mesmos valores t_i dados pelo simulador IMEX, como mostrado na Tabela 8. A Figura 40 mostra a curva obtida aproximada pelo modelo segmentado e a curva original de produção de óleo fornecida pelo simulador IMEX.

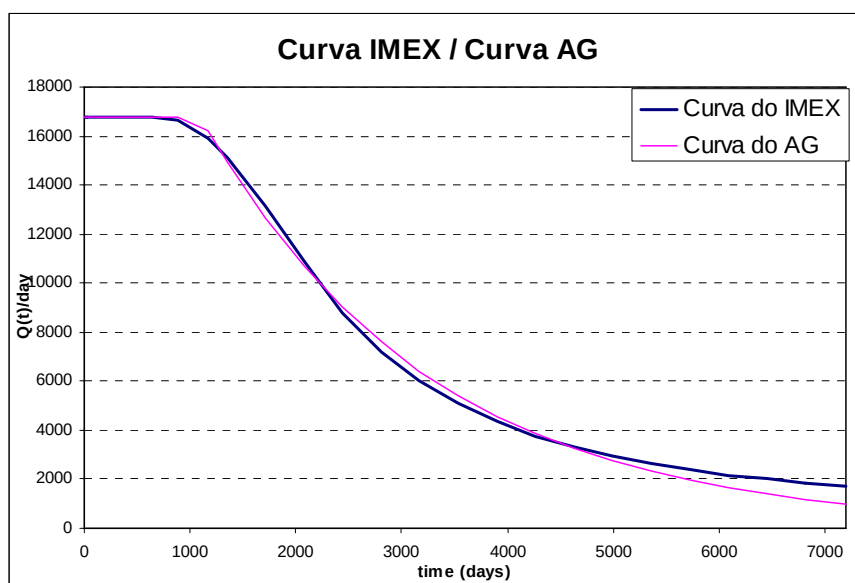


Figura 40. Curva obtida e curva original do Exp-1: 4 *five-spots*

5.3.1.2.

Experimento 02: alternativa aleatória

Neste experimento, a alternativa utilizada foi escolhida a partir da geração aleatória do gerador de alternativas empregado para criar as amostras para o treinamento dos módulos de aproximação. A Figura 41 mostra a curva de óleo acumulado e a Figura 42 mostra a curva produção diária de óleo, ambas obtidas através do simulador de reservatórios.

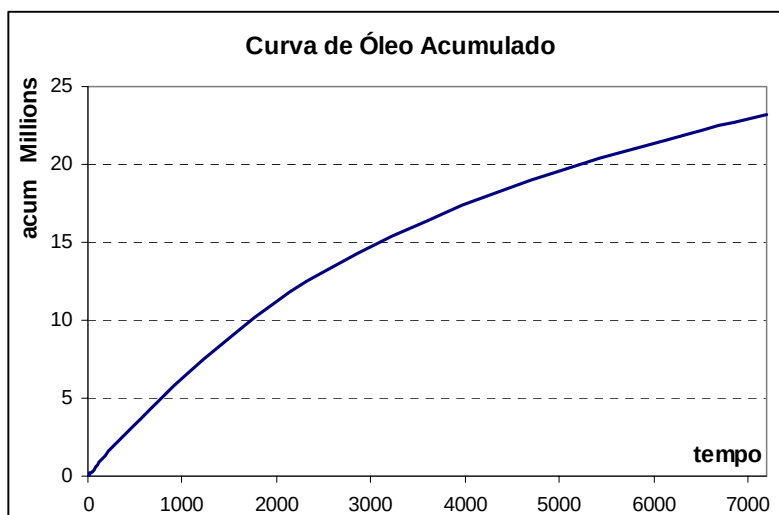


Figura 41. Curva de óleo acumulado: alternativa aleatória

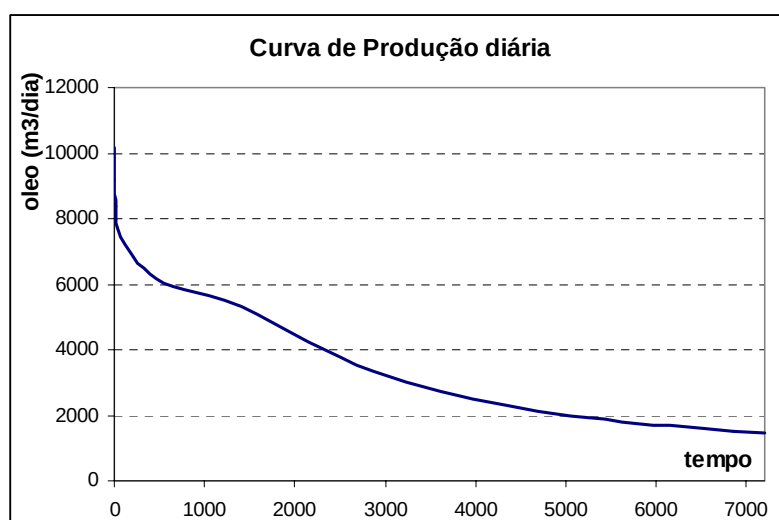


Figura 42. Curva de produção de óleo: alternativa aleatória

Para realizar a execução do algoritmo genético foram empregados os mesmos parâmetros da Tabela 6 e as mesmas taxas da Tabela 7 que foram usados no experimento 1. Após a otimização dada pelo algoritmo genético, foram obtidos os resultados que se mostram na Tabela 9, a seguir.

t_a	t_b	t_c	q_a	q_b	ρ	MAPE
8.638	19.17	41.08	630.53	7346.05	0.097	2.38%

Tabela 9. Resultados obtidos para o Exp- 2: alternativa aleatória.

Como pode se ver na Tabela 9, o erro MAPE de aproximação obtido é de 2.38%. A curva obtida aplicando a função segmentada da equação (23) para os parâmetros obtidos pelo AG para a alternativa aleatória é mostrada na Figura 43, a seguir.

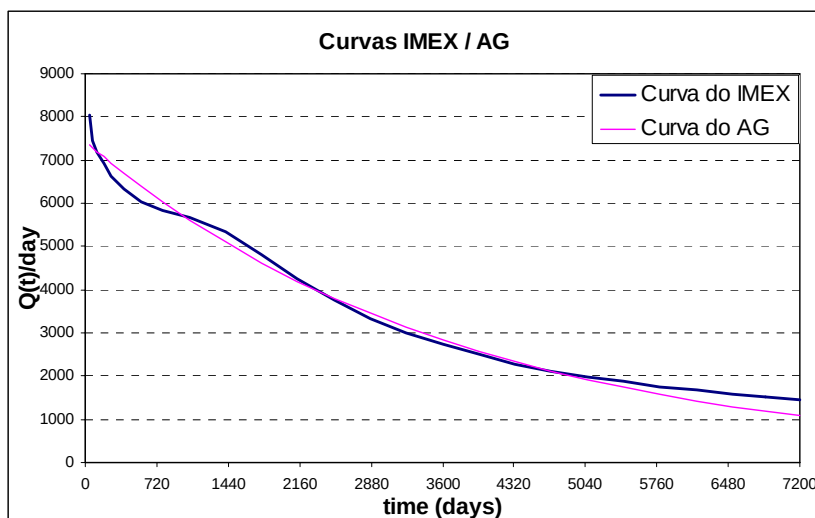


Figura 43. Curva obtida e curva original do Exp-2: alternativa aleatória

5.3.1.3.

Experimento 3: cronograma de poços

Neste experimento foi usada a configuração 5 *five-spots* do experimento 1. Neste caso, a configuração contém cronograma de poços, isto significa que os poços não começam a produzir todos juntos no tempo zero (t_0), mas em diferentes tempos e com deslocamentos de tempo de 60 dias. Cabe ressaltar que a ordem de abertura dos poços foi obtida a partir de um sistema de otimização de cronograma de poços baseado em algoritmos genéticos. A Figura 44 mostra a curva de óleo acumulado e a Figura 1 mostra a curva de produção diária de óleo para esta alternativa considerando a entrada paulatina de poços no tempo.

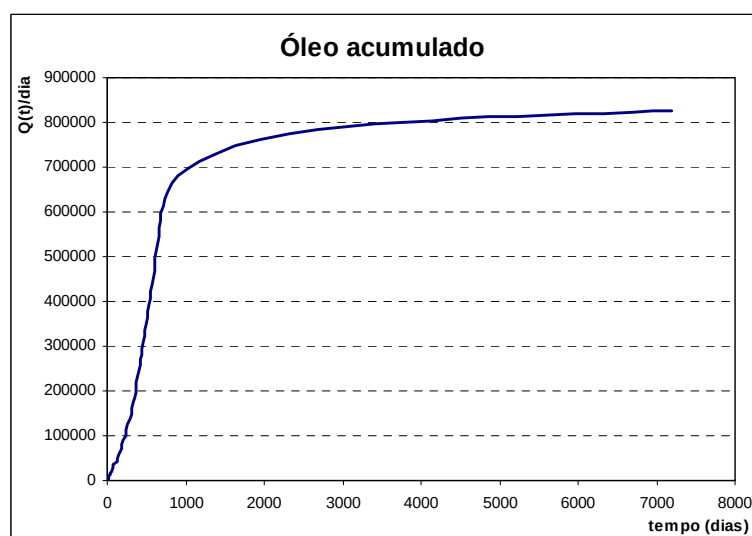


Figura 44. Curva de óleo acumulado: cronograma de poços

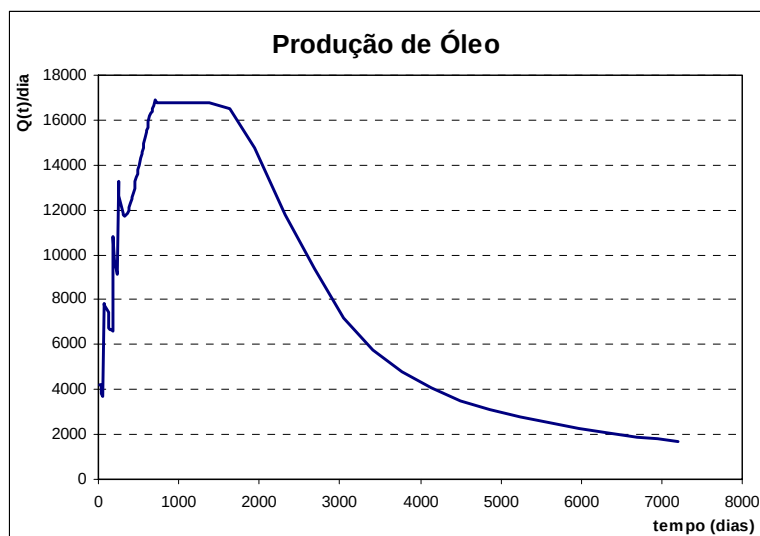


Figura 45. Curva de produção diária de óleo: cronograma de poços

Pode ser visto na Figura 45 a tendência ascendente ocorrida entre o dia 0 e o dia 660 devida à entrada paulatina dos poços.

O algoritmo genético foi executado empregando os mesmos parâmetros da Tabela 6 e taxas da Tabela 7. Após a execução, os resultados obtidos mostraram um incremento no erro de aproximação MAPE, como pode se visto na Tabela 10 a seguir.

t_a	t_b	t_c	q_a	q_b	ρ	MAPE
70,87	665	1703,55	7130,34	16691	0.2080	2.06%

Tabela 10. Resultados obtidos para o Exp-3: cronograma de poços.

O erro de aproximação MAPE foi de 2.06%. Na Figura 46, pode se visualizar como a curva segmentada acompanha o formato mais complicado da curva obtida por simulação.

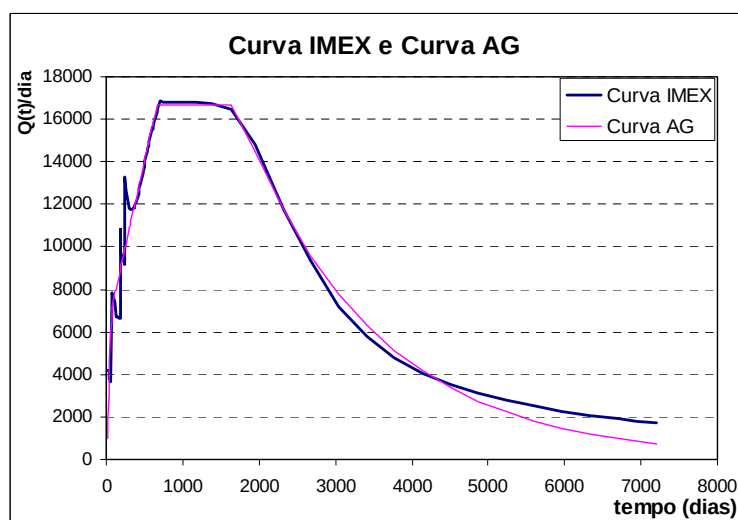


Figura 46. Resultados gráficos do Exp-3: cronograma de poços

5.3.2. Geração dos padrões de aprendizado

Este estágio dos experimentos de aproximação de curvas de produção consiste em obter a lista de padrões de entrada e saída, onde cada um destes padrões representa uma alternativa que contém, como entrada, uma lista de poços produtores e injetores localizados aleatoriamente no campo petrolífero, e os respectivos pontos definidos da curva de produção, no caso do modelo de curva de produção, ou os respectivos parâmetros já obtidos através da aproximação AG descrita na seção (4.5.3), no caso do modelo de curva segmentada. Nas seções (4.5.4.1 e 4.5.4.2) foram mostrados os formatos dos padrões de entrada e saída para o modelo de curva de produção e modelo de curva segmentada respectivamente.

Durante a geração dos padrões foram levadas em consideração as seguintes regras:

- Não gerar configurações sem poços produtores;
- Não podem haver poços repetidos, isto é, dois ou mais poços na mesma locação, dentro de um mesmo padrão;
- Foram utilizados geradores de números quase-aleatórios (Niederreiter, 1992) de Sobol (Sobol, 1967) durante a geração dos padrões de treinamento, isto visando obter um bom espalhamento das amostras no espaço de busca, como é característico das seqüências de números quase aleatórios;
- Utilização de geradores de números pseudo-aleatórios para a criação dos padrões de teste e validação.

Um diagrama de blocos que mostra a seqüência de geração dos padrões mostra-se na Figura 47 a seguir.

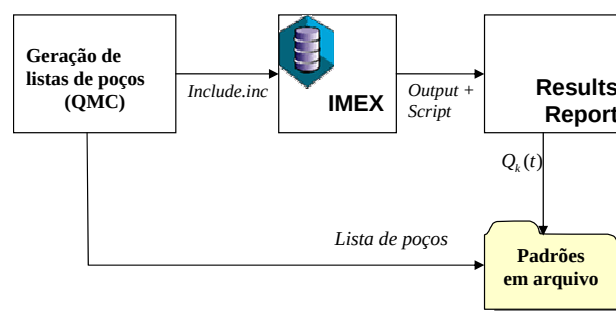


Figura 47. Obtenção de padrões de entrada e saída

Como pode ser visto na Figura 47, primeiramente é gerada uma lista de poços produtores e injetores conforme as regras descritas acima. A lista de poços é armazenada localmente e enviada para ser simulada. Após a execução do simulador IMEX e o programa *results report* é obtido um arquivo texto .RWO contendo todos os pontos da curva de produção de óleo gerados pela simulação.

Para o modelo de pontos de curva de produção, a partir do arquivo .RWO são extraídos os pontos escolhidos, isto é, os valores de óleo acumulado correspondentes aos tempos t_{90} , t_{180} , t_{360} , t_{720} , t_{1800} , t_{3600} , t_{5400} e t_{7200} .

Para o modelo de curva segmentada, do arquivo .RWO são carregados os valores da curva de produção que será utilizada como referência na função objetivo do algoritmo genético, em seguida executa-se o algoritmo genético estimando os 06 parâmetros da curva segmentada, vindo a ser estes, a saída correspondente à lista de poços.

Para qualquer um dos modelos, é montada o padrão de entrada e saída juntando a lista de poços e os valores obtidos após a simulação e, finalmente este padrão completo é armazenado localmente para, ao final da geração e simulação de todos os padrões solicitados, estes sejam armazenados em arquivo de texto para seu posterior uso durante os testes de aprendizado dos modelos aproximadores.

Cabe ressaltar que, para a geração de padrões de treinamento, também é possível utilizar o *framework* de avaliações distribuídas desenvolvido neste trabalho.

5.3.3. Experimentos de aprendizado dos modelos aproximadores

Estes experimentos foram realizados com o intuito de encontrar as arquiteturas de redes neurais e parametrizações mais adequadas ao problema. Os experimentos realizados envolveram os dois modelos de aproximação propostos.

Para realizar estes experimentos, foram gerados 4000 padrões com configurações de 2 poços (injetor e produtor), 10000 padrões para configurações de 8 poços (4 produtores e 4 injetores) e 20.000 padrões para 10 poços (5 produtores e 5 injetores). Os valores de saída foram os 8 pontos da curva de produção acumulada de óleo, gás e água e os 6 parâmetros da curva segmentada obtidos pela otimização AG.

Nas seções seguintes (5.3.3.1 até 5.3.3.4) mostram-se os experimentos realizados para o modelo de curva de óleo acumulado para redes *feed-forward* com aprendizado *back-propagation*, redes recorrentes de Elman, modelo NFHB e redes RBF, para configurações com 2, 8 e 10 poços.

5.3.3.1.

Redes *feed-forward* – Pontos de Curva

Para realizar este experimento, foi escolhido o uso de taxa adaptativa de aprendizado e termo de momento para o treinamento e usou-se o gradiente descendente para a otimização dos pesos. A arquitetura que forneceu melhores resultados foi uma *feed-forward* de $2n$ entradas onde n é o número de poços; com 35 neurônios na camada escondida e as 8 saídas para os valores de óleo em tempos 90, 180, 360, 720, 1800, 3600, 5400 e 7200 dias, como definido no modelo de pontos de curvas de óleo acumulado (vide seção 4.5.2.1). O processo de aprendizado foi feito com padrões para 2 poços, padrões para 8 poços e padrões para 10 poços.

Os padrões de entrada-saída foram gerados para o modelo de reservatório heterogêneo da Figura 36. Para todos os casos, os padrões foram distribuídos da seguinte forma:

70% para o treinamento, 20% para validação e 10% para testes

Os parâmetros de aprendizado iniciais foram os que mostra a Tabela 11 a seguir

Parâmetro de Aprendizado	Valor
taxa de aprendizado inicial	0.1
Termo de momento	0.9
multiplicador para incremento de taxa	1.1
multiplicador para decremento de taxa	0.7

Tabela 11. Parâmetros de aprendizado – Rede *feed-forward*

Foi utilizado o método de validação cruzada (Stone, 1974; 1978) para escolha da melhor arquitetura e para determinar o fim do treinamento.

As tabelas Tabela 12, Tabela 13 e Tabela 14 mostram os resultados obtidos com o algoritmo de aprendizado *back-propagation* usando taxa adaptativa e termo de momento para padrões com 2, 8 e 10 poços.

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	2800	800	400	1.98%	2.01%	2.04%
180	2800	800	400	2.25%	2.27%	2.29%
360	2800	800	400	2.79%	2.86%	2.88%
720	2800	800	400	2.70%	2.72%	2.75%
1800	2800	800	400	2.28%	2.30%	2.35%
3600	2800	800	400	2.66%	2.70%	2.75%
5400	2800	800	400	3.49%	3.55%	3.70%
7200	2800	800	400	3.45%	3.55%	3.63%

Tabela 12. Resultados da rede *feed-forward* – 2 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	7000	2000	1000	4.84%	4.88%	5.07%
180	7000	2000	1000	4.65%	4.66%	4.73%
360	7000	2000	1000	4.35%	4.38%	4.43%
720	7000	2000	1000	4.19%	4.25%	4.29%
1800	7000	2000	1000	4.27%	4.29%	4.36%
3600	7000	2000	1000	4.75%	4.75%	4.85%
5400	7000	2000	1000	4.80%	4.86%	4.90%
7200	7000	2000	1000	4.87%	4.89%	4.91%

Tabela 13. Resultados da rede *feed-forward* – 8 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	14000	4000	2000	5.19%	5.34%	5.36%
180	14000	4000	2000	4.86%	5.03%	5.14%
360	14000	4000	2000	4.75%	4.78%	4.84%
720	14000	4000	2000	4.63%	4.71%	4.73%
1800	14000	4000	2000	5.08%	5.16%	5.19%
3600	14000	4000	2000	5.67%	5.75%	5.80%
5400	14000	4000	2000	5.82%	5.84%	5.87%
7200	14000	4000	2000	5.87%	5.88%	5.90%

Tabela 14. Resultados da rede *feed-forward* – 10 poços.

5.3.3.2.

Redes de Elman – Pontos de Curva.

As redes de Elman foram configuradas de forma similar às redes *feed-forward*, isto é, utilizando o gradiente descendente para a otimização dos pesos e algoritmo *back-propagation* com taxa adaptativa de aprendizado e termo momento para o aprendizado. O método de validação cruzada também foi empregado para escolher a melhor arquitetura e o tempo de finalização do treinamento. Os padrões foram divididos em 70% para o treinamento, 20% para validação e 10% para testes

Os parâmetros de aprendizado iniciais empregados foram os mostrados pela Tabela 15 a seguir.

Parâmetro de Aprendizizado	Valor
taxa de aprendizado inicial	0.1
Termo de momento	0.9
multiplicador para incremento de taxa	1.1
multiplicador para decremento de taxa	0.7

Tabela 15. Parâmetros de aprendizado – Rede de Elman

Para os experimentos realizados com redes de Elman, a melhor arquitetura obtida contém 30 neurônios na camada escondida.

As Tabela 16, Tabela 17 e Tabela 18 contêm os resultados obtidos com o algoritmo de aprendizado *back-propagation* usando taxa adaptativa e termo de momento para padrões com 2, 8 e 10 poços, como se mostra a seguir.

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	2800	800	400	1.80%	1.82%	1.89%
180	2800	800	400	1.85%	1.87%	1.89%
360	2800	800	400	1.88%	1.91%	1.92%
720	2800	800	400	1.90%	1.92%	1.93%
1800	2800	800	400	2.01%	2.08%	2.10%
3600	2800	800	400	2.34%	2.40%	2.41%
5400	2800	800	400	2.45%	2.47%	2.51%
7200	2800	800	400	2.56%	2.58%	2.65%

Tabela 16. Resultados da rede Elman – 2 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	7000	2000	1000	4.82%	4.84%	4.90%
180	7000	2000	1000	4.59%	4.62%	4.67%
360	7000	2000	1000	4.89%	4.92%	4.93%
720	7000	2000	1000	4.95%	4.98%	4.99%
1800	7000	2000	1000	4.70%	4.73%	4.78%
3600	7000	2000	1000	4.67%	4.70%	4.71%
5400	7000	2000	1000	4.56%	4.61%	4.67%
7200	7000	2000	1000	4.78%	4.81%	4.86%

Tabela 17. Resultados da rede Elman – 8 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	14000	4000	2000	5.50%	5.64%	5.77%
180	14000	4000	2000	5.30%	5.40%	5.67%
360	14000	4000	2000	5.07%	5.21%	5.41%
720	14000	4000	2000	5.51%	5.76%	5.86%
1800	14000	4000	2000	5.72%	5.80%	5.98%
3600	14000	4000	2000	5.66%	5.71%	5.75%
5400	14000	4000	2000	5.15%	5.52%	5.64%
7200	14000	4000	2000	5.25%	5.29%	5.35%

Tabela 18. Resultados da rede Elman – 10 poços

5.3.3.3.

Modelos NFHB – Pontos de Curva

Para realizar os experimentos com o modelo NFHB, foi alterado o modelo NFHB existente de modo a não empregar seleção de variáveis, isto é, escolher dentre as variáveis de entrada àquelas que contribuem com mais informação para a saída, e descartar aquelas que contribuem com menos informação; descartar alguma variável implica em desconsiderar algum parâmetro de poço das alternativas presentes nos padrões de entrada-saída gerados, perdendo-se o sentido do aprendizado.

Os experimentos com o modelo NFHB foram realizados para o modelo de pontos de curva de produção variando o parâmetro de taxa de decomposição. O valor de taxa de decomposição que permitiu obter melhores aproximações foi o valor 0.007 para todos os casos.

Nas Tabela 19, Tabela 20 e Tabela 21 são mostrados os resultados dos experimentos realizados com o modelo NFHB, e seguem a continuação.

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	1400	400	200	0.98%	0.99%	0.99%
180	1400	400	200	0.95%	0.97%	0.99%
360	1400	400	200	0.96%	0.97%	0.99%
720	1400	400	200	0.77%	0.81%	0.88%
1800	1400	400	200	0.81%	0.82%	0.84%
3600	1400	400	200	0.72%	0.74%	0.75%
5400	1400	400	200	0.80%	0.81%	0.82%
7200	1400	400	200	0.56%	0.58%	0.65%

Tabela 19. Resultados do modelo NFHB – 2 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	7000	2000	1000	4.52%	4.56%	4.61%
180	7000	2000	1000	4.75%	4.78%	4.82%
360	7000	2000	1000	4.75%	4.75%	4.78%
720	7000	2000	1000	4.63%	4.69%	4.70%
1800	7000	2000	1000	3.99%	4.05%	4.06%
3600	7000	2000	1000	4.12%	4.15%	4.19%
5400	7000	2000	1000	4.01%	4.03%	4.09%
7200	7000	2000	1000	3.98%	4.01%	4.08%

Tabela 20. Resultados do modelo NFHB – 8 poços

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	14000	4000	2000	4.39%	4.40%	4.43%
180	14000	4000	2000	4.35%	4.36%	4.40%
360	14000	4000	2000	4.35%	4.40%	4.45%
720	14000	4000	2000	4.38%	4.44%	4.51%
1800	14000	4000	2000	4.45%	4.50%	4.55%
3600	14000	4000	2000	4.34%	4.40%	4.42%
5400	14000	4000	2000	4.15%	4.18%	4.20%
7200	14000	4000	2000	4.12%	4.15%	4.17%

Tabela 21. Resultados do modelo NFHB – 10 poços

5.3.3.4. Redes RBF – Pontos de Curva

Nos experimentos realizados com as redes RBF, o parâmetro que foi variado a fim de determinar a melhor arquitetura foi o '*spread*' que significa a variância inicial ou 'cobertura' de cada um dos neurônios da camada escondida que representam funções *gaussianas* sendo inseridos na arquitetura durante o processo de treinamento conforme a seção (3.2.4.2). A Tabela 22 mostra os melhores resultados obtidos com as redes RBF, a seguir.

Saída (ti)	Trein.	Valid	Teste	MAPE Trein	MAPE Valid	MAPE Teste
90	14000	4000	2000	0%	5.37%	5.45%
180	14000	4000	2000	0%	4.52%	4.60%
360	14000	4000	2000	0%	3.54%	3.50%
720	14000	4000	2000	0%	3.34%	3.43%
1800	14000	4000	2000	0%	2.54%	2.75%
3600	14000	4000	2000	0%	1.87%	1.98%
5400	14000	4000	2000	0%	1.66%	1.78%
7200	14000	4000	2000	0%	1.56%	1.58%

Tabela 22. Resultados do modelo RBF – 2 poços

O erro de treinamento que aparece na Tabela 22 como 0% não é significativo dada a metodologia de aprendizado das redes RBF, na qual, as funções radiais reproduzem perfeitamente os valores utilizados no treinamento.

Não foram realizados testes para 8 e 10 poços por limitações computacionais dado o número entradas e saídas dos padrões.

Nas seções seguintes (5.3.3.5 até 5.3.3.7) mostram-se os experimentos realizados com o modelo de curva segmentada para redes *feed-forward* com aprendizado *back-propagation*, redes recorrentes de Elman e modelo NFHB, para configurações com 10 poços.

5.3.3.5. Redes *feed-forward* – Curva Segmentada

Para realizar este experimento, foi empregado o algoritmo de aprendizado *back-propagation* com taxa de aprendizado adaptativa e termo de momento. Para a otimização dos pesos empregou-se o gradiente descendente. Para escolher a melhor arquitetura e determinar o fim do processo de treinamento foi usado o método de validação cruzada. A arquitetura que forneceu melhores resultados foi uma *feed-forward* de 20 entradas que significa até 10 poços, com 35 neurônios na camada escondida e 6 saídas, uma para cada parâmetro do modelo segmentado. O processo de aprendizado foi feito com 4000 padrões, dentre os quais foram usados 2800 para treinamento, 800 para validação e 400 padrões para testes. A taxa de aprendizado inicial foi de 0.1 e termo de momento 0.9. Foi empregado o método de validação cruzada para escolha da melhor arquitetura e determinar o fim do treinamento.

A Tabela 23 mostra os resultados obtidos com redes *feed-forward* e algoritmo de aprendizado *back-propagation*, como segue a continuação.

Rede BP	Numero de Padrões			Erros MAPE		
	Treino	Valid	Teste	Treino	Valid	Teste
Ta	2800	800	400	0.09%	0.12%	0.13%
Tb	2800	800	400	0.10%	0.12%	0.16%
Tc	2800	800	400	4.45%	4.51%	4.47%
Qa	2800	800	400	4.45%	4.78%	4.87%
Qb	2800	800	400	5.03%	5.08%	5.45%
Rho	2800	800	400	5.23%	5.31%	5.59%

Tabela 23. Resultados da rede *feed-forward* – Curva segmentada

5.3.3.6. Redes de Elman – Curva segmentada

Os experimentos com a arquitetura de rede recorrente de Elman, também empregaram o algoritmo de aprendizado *back-propagation* com taxa adaptativa e termo de momento, de forma similar às redes *feed-forward* do experimento anterior. Para escolher a melhor arquitetura e determinar a finalização do treinamento foi empregada a metodologia de validação cruzada.

A arquitetura que forneceu melhores resultados foi uma rede com 20 neurônios na entrada, 50 neurônios na camada escondida e 6 neurônios na camada de saída.

Os resultados após o treinamento, validação e testes se mostram na Tabela 24 a seguir

Rede Elman	Número de Padrões			Erros (MAPE)		
	Treino	Valid	Teste	Treino	Valid	Teste
Ta	2800	800	400	0.08%	0.09%	0.10%
Tb	2800	800	400	0.09%	0.11%	0.12%
Tc	2800	800	400	4.45%	4.51%	4.47%
Qa	2800	800	400	4.45%	4.78%	4.87%
Qb	2800	800	400	5.00%	5.05%	5.22%
Rho	2800	800	400	5.21%	5.26%	6.09%

Tabela 24. Resultados da rede Elman – Curva segmentada

Cabe ressaltar que as redes de Elman alcançaram o melhor valor de treinamento antes de chegar ao estado de *overfitting* em menor número de iterações do que as redes *feed forward*. Isto se reflete em menor tempo computacional requerido para o treinamento

5.3.3.7.

Modelos *Neuro-Fuzzy* – Curva segmentada

Os experimentos que empregaram o modelo *neuro-fuzzy* hierárquico para aproximar a curva de produção utilizando o modelo de curva segmentada foram realizados variando o parâmetro taxa de decomposição. Este parâmetro permite controlar o crescimento da estrutura de conseqüentes evitando o crescimento excessivo desta estrutura o que na prática cria problemas de generalização, como mencionado na seção (3.3.2.3) e também de exigências computacionais.

Após realizados vários experimentos visando ajustar os parâmetros do modelo NFHB, encontrou-se que o valor da taxa de decomposição que permite o melhor equilíbrio entre tamanho de estrutura, generalização e boa aproximação para o problema em questão é o valor 0.007.

A Tabela 25 apresenta resultados dos experimentos realizados com o modelo *neuro-fuzzy* hierárquico já com a taxa de decomposição 0.007, como se mostra a seguir

Modelo NFHB	Número de Padrões			Erros (MAPE)		
	Treino	Valid	Testes	Treino	Valid	Testes
Ta	2800	800	400	0.05%	0.06%	0.05%
Tb	2800	800	400	0.08%	0.09%	0.09%
Tc	2800	800	400	3.25%	3.31%	3.37%
Qa	2800	800	400	3.36%	3.42%	3.47%
Qb	2800	800	400	3.33%	3.38%	3.61%
Rho	2800	800	400	4.23%	4.24%	4.49%

Tabela 25. Resultados do modelo NFHB – Curva segmentada

É interessante ressaltar que o modelo *neuro fuzzy* hierárquico NFHB utilizado requer menor tempo para realizar o processo de aprendizado, se comparado com os tempos requeridos no treinamento de redes neurais. Isto não desconta o investimento computacional requerido para obter as amostras de aprendizado.

5.4. Experimentos com o Sistema de Otimização

Nesta seção são mostrados os experimentos realizados com o sistema de otimização desenvolvido neste trabalho, isto é, experimentos que fornecem como resposta, uma alternativa de desenvolvimento do campo petrolífero.

As realizações destes experimentos envolvem os seguintes pontos:

- a execução do algoritmo genético;
- a utilização do ambiente paralelo nas avaliações distribuídas;
- a geração do mapa de qualidade e uso das informações fornecidas por ele.

Para a realização dos experimentos de otimização é necessário definir os valores dos parâmetros necessários para o cálculo do valor presente líquido, estes parâmetros são divididos em parâmetros relacionados ao mercado, parâmetros do investimento (*Capital expenditure* – CAPEX) e relacionados ao custo operacional (*Operation Expenditure* – OPEX). As Tabela 26, Tabela 27 e Tabela 28 mostram os parâmetros empregados nos experimentos realizados neste trabalho.

Variáveis de mercado	Valor
Preço do petróleo (US\$)	20.0
Royalties (Ry)	0.1
Taxas sociais + imposto renda (I)	0.34
Taxa de desconto	0.1

Tabela 26. Parâmetros relacionados às variáveis de mercado.

Custos para desenvolvimento (CAPEX)	Valor
Perfuração + árvore de natal (US\$)	20.000.000
Riser (US\$)	2.000.000
Lines bundle (US\$/km)	2.000.000
Planta (US\$)	50.000.000
Plataforma (US\$)	350.000.000
Fator de custo poço horizontal	50%

Tabela 27. Parâmetros CAPEX.

Custos operacionais (OPEX)	Valor
Custos Fixos	0
Custos Variáveis	5
Manutenção de poço (US\$/ano)	600.000
Custo retirada d'água (US\$/m ³)	1

Tabela 28. Parâmetros OPEX.

Nas seções seguintes, mostram-se os experimentos realizados com o campo homogêneo 30x30x1 e experimentos realizados com o campo heterogêneo.

5.4.1. Experimentos com reservatório 30x30x1

Os experimentos realizados com o reservatório 30x30x1 homogêneo permitiram ajustar os parâmetros do algoritmo genético. Os melhores resultados foram obtidos com os valores que mostra a Tabela 29, a seguir.

Parâmetros do AG	Valor
Número de gerações	200
Número de indivíduos	120
Número de rodadas	1
<i>Steady-State</i>	0.4
Probabilidade de criar poço	0.5
Taxa de <i>crossover</i> inicial	0.65
Taxa de mutação inicial	0.08

Tabela 29. Parâmetros do AG: reservatório 30x30x1

Na realização dos experimentos, foram consideradas as restrições definidas na seção (4.3.3), os valores se mostram na Tabela 30, a seguir

Restrições no AG	Valor(m)
Distância mínima entre poços	400
Máximo comprimento de poço horizontal	1500

Tabela 30. Restrições nas soluções do AG: reservatório 30x30x1

Com estes parâmetros e restrições já fixados, nas seções seguintes se mostram 3 experimentos realizados.

5.4.1.1. Experimento 1 – Inicialização Aleatória

Neste experimento, foram evoluídas soluções, partindo da inicialização aleatória, usando poços horizontais e verticais e as restrições da Tabela 30. As curvas de evolução obtidas se mostram na Figura 48, a seguir.

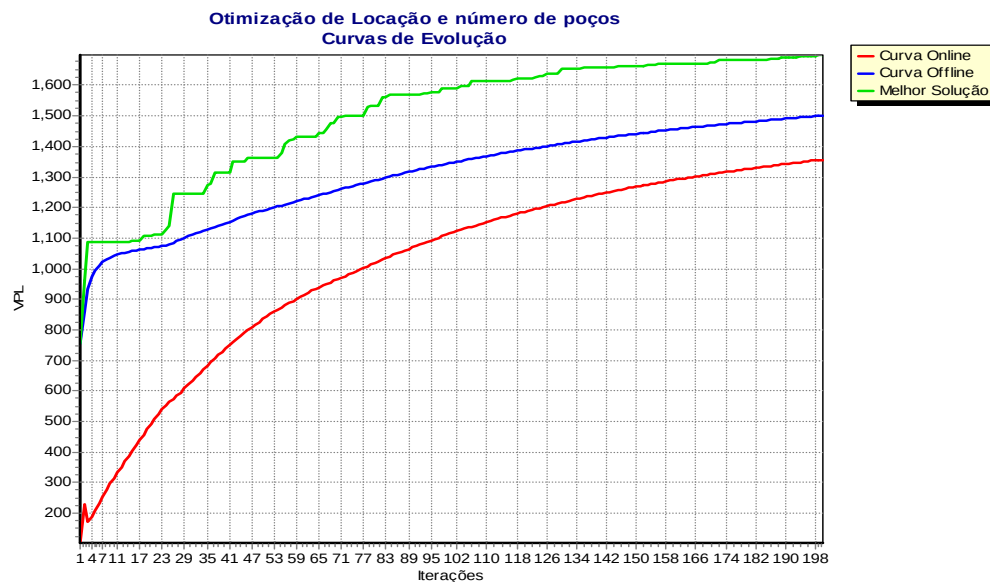


Figura 48. Curvas de evolução: reservatório 30x30x1 – Exp 1.

A melhor solução (cromossoma) obtida neste experimento tem um valor de avaliação de

VPL	1.699.122.384,65 (US\$)
óleo recuperado	364.075.278 (bbl)

Na Figura 49, mostra-se a disposição de poços obtida para esta alternativa. Na Figura 50 mostra-se a curva de óleo acumulado respectiva.

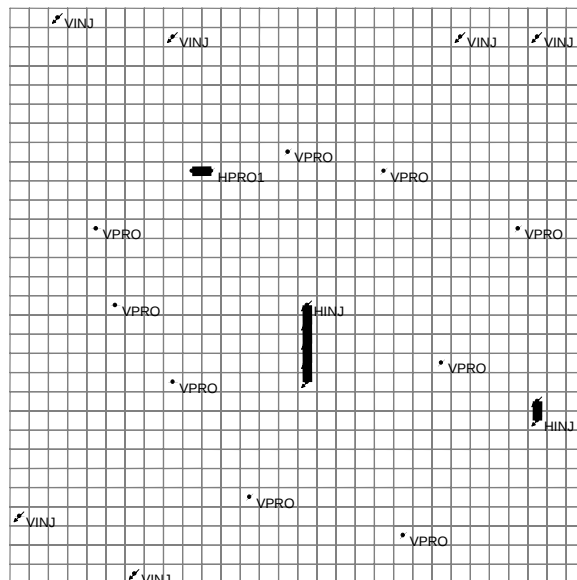


Figura 49. Disposição dos poços: reservatório 30x30x1 – Exp 1,

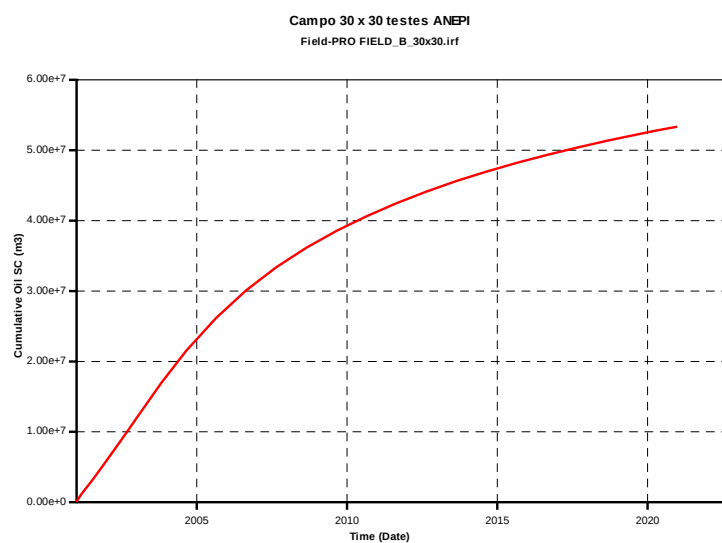


Figura 50. Curva de óleo acumulado: reservatório 30x30x1 – Exp 1,

Cabe ressaltar que este experimento foi realizado utilizando o algoritmo genético com avaliações distribuídas. O ambiente paralelo empregado possuiu 07 computadores da rede do laboratório ICA. A Tabela 31 contém a lista de computadores utilizados neste experimento, onde é detalhado o nome do microcomputador, a velocidade e tipo de processador, como se mostra a seguir:

Nome do computador	Velocidade
ICA – LAMBADA	AMD-XP2800
ICA – BOLERO	INTEL-P3 1G
ICA – SWING	INTEL-P2 400M
ICA – ROCK	AMD-XP2200
ICA – HEAVYMETAL	AMD-XP2200
ICA – TUNTUNA	INTEL-P2 300M
ICA – HUAYLAS	AMD-XP2400

Tabela 31. Computadores utilizados: reservatório 30x30x1 – Exp 1

O tempo de execução das 8000 avaliações da otimização realizada neste experimento foi 1.20 horas. Se esta otimização fosse realizada em um único computador AMD-XP-2800 que é o mais rápido da lista, o tempo de execução aproximado seria de aproximadamente 5.00 horas. Neste caso, o custo de comunicação e manuseio de processos paralelos não é muito desprezível se comparado com o tempo que a avaliação em si leva.

Na Figura 51 a seguir, mostra-se um resumo do uso dos diferentes computadores após finalizadas as 8000 avaliações.

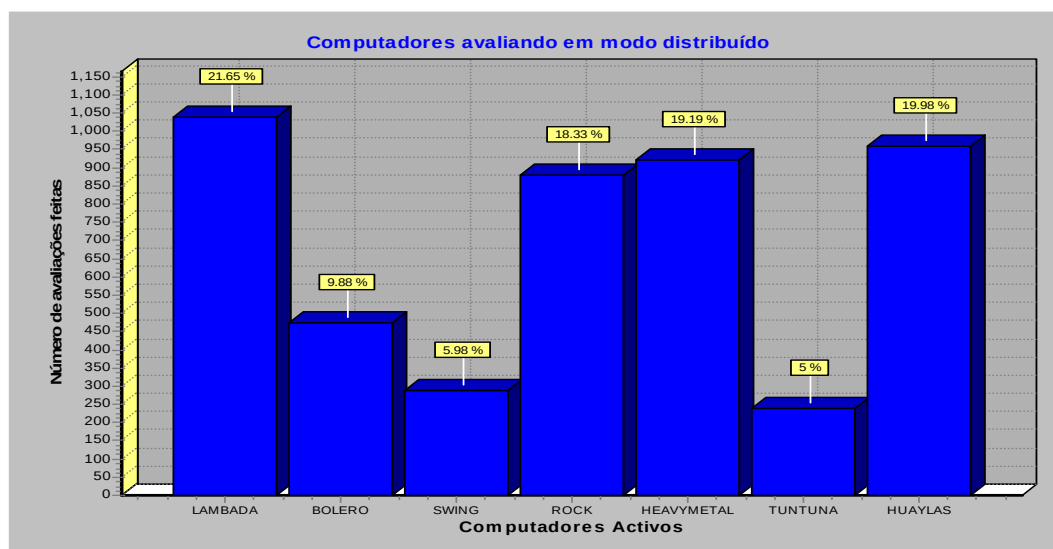


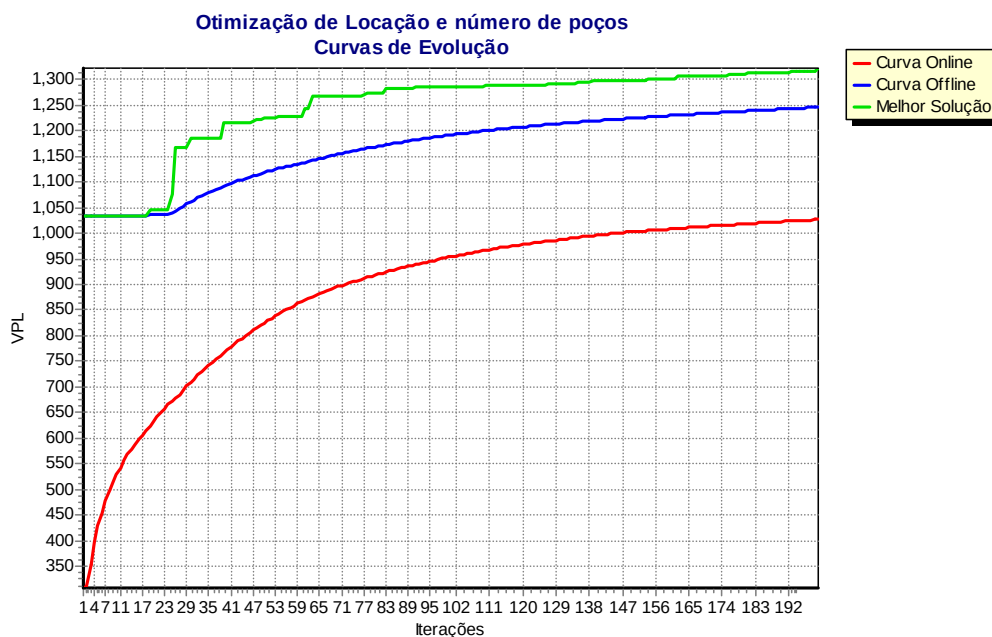
Figura 51. Percentual de avaliação dos computadores – Exp 1.

Na figura anterior, cada barra representa o percentual de avaliações que realizou cada computador da lista, o conjunto de barras permite visualizar graficamente uma comparação de quanto aportou cada um dos computadores no total de avaliações realizadas durante a otimização.

5.4.1.2. Experimento 2 – Poços Verticais

Neste experimento, foram evoluídas soluções permitindo ao algoritmo genético usar apenas poços verticais. Para este caso mantém-se a restrição de distância mínima de 400 m da Tabela 30.

As curvas de evolução obtidas: *on-line* e *off-line* definidas na seção (3.4.6) e a curva da melhor solução encontrada, se mostram na Figura 52 a seguir



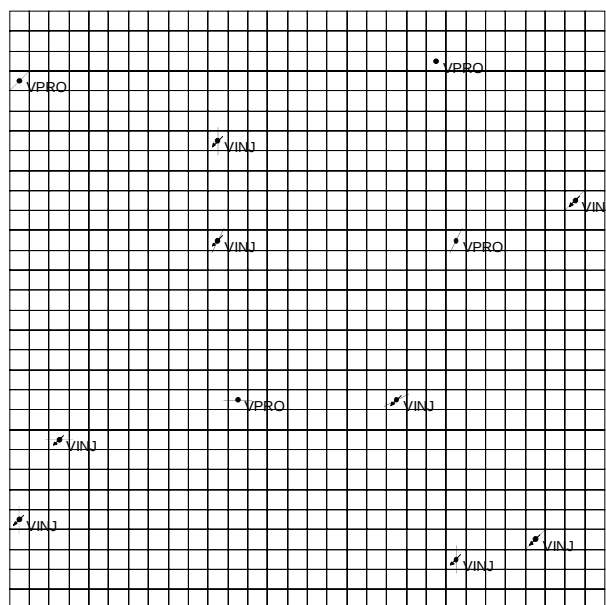


Figura 53. Disposição dos poços: reservatório 30x30x1 – Exp 2.

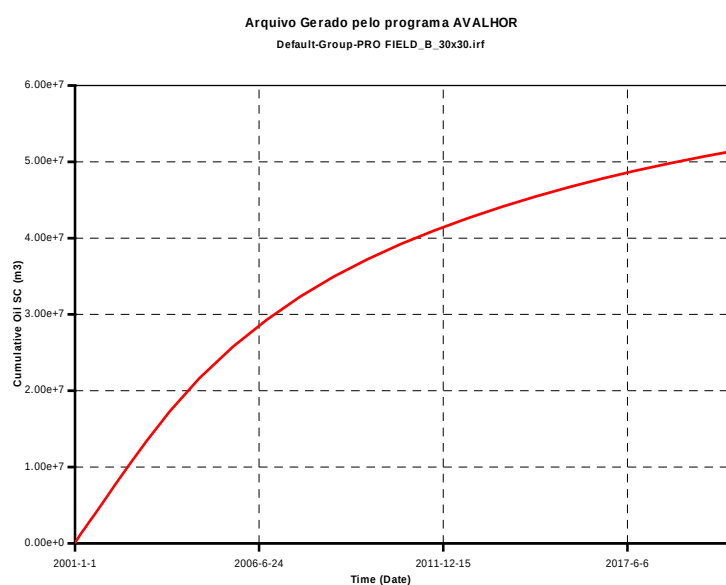


Figura 54. Curva de óleo acumulado: reservatório 30x30x1 – Exp 2.

Na alternativa, obtida tanto o valor do VPL quanto o valor de óleo acumulado foram menores que os obtidos no experimento 1, isto se deve à limitação imposta para evoluir soluções empregando apenas poços verticais.

5.4.1.3.

Experimento 3 – Semente Inicial

Este experimento evoluiu soluções com poços verticais e horizontais no campo homogêneo, para sua realização, foi inserida na população inicial do algoritmo genético uma semente que consiste em uma configuração clássica 4-*five-spots*. O VPL para a configuração é US\$ 1.397.981.156,05 com uma extração de óleo de 297.025.122 bbl. O algoritmo foi executado com os seguintes parâmetros:

Parâmetros do AG	Valor
Número de gerações	200
Número de indivíduos	60
Número de rodadas	1
<i>Steady-State</i>	0.4
Probabilidade de poço	0.5
Taxa de <i>crossover</i> inicial	0.65
Taxa de mutação inicial	0.08

Tabela 32. Parâmetros do AG: reservatório 30x30x1 – semente inicial

As curvas de evolução *offline*, *online* e de melhor elemento, obtidas neste experimento se mostram na Figura 55 a seguir.

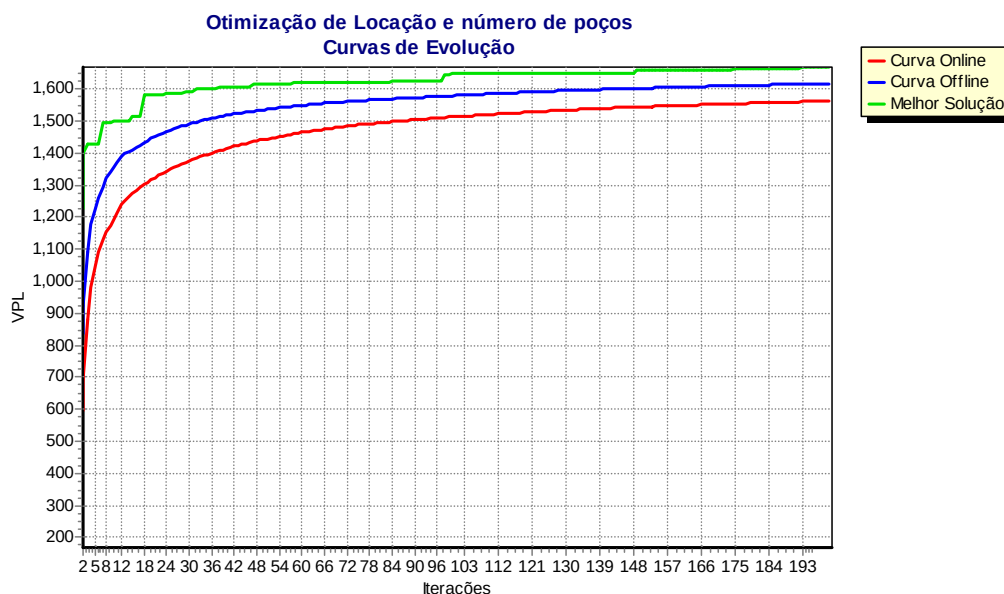


Figura 55. Curvas de evolução: reservatório 30x30x1 – Exp 3.

Na Figura 55, se mostra que já nas primeiras gerações o algoritmo genético encontrou soluções com melhores valores de VPL do que a semente inicial. Na segunda geração, o VPL do melhor elemento foi de US\$1.429.235.042,91.

Após a execução das 200 gerações, a melhor solução encontrada teve os seguintes valores de avaliação.

VPL **1.667.820.970,00 (US\$)**

óleo recuperado **359.100.000 (bbl)**

Estes valores são superiores aos que foram obtidos na avaliação da semente inicial. Na Figura 56 se mostra a disposição de poços da configuração da semente inicial *4-five-spots* e, na Figura 57 a disposição da configuração encontrada pelo algoritmo genético após as 200 gerações.

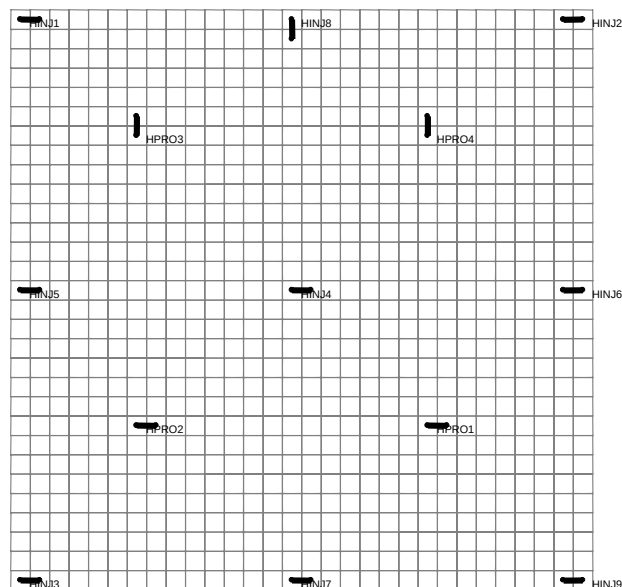


Figura 56. Disposição dos poços do 4-five spots inicial – Exp 3

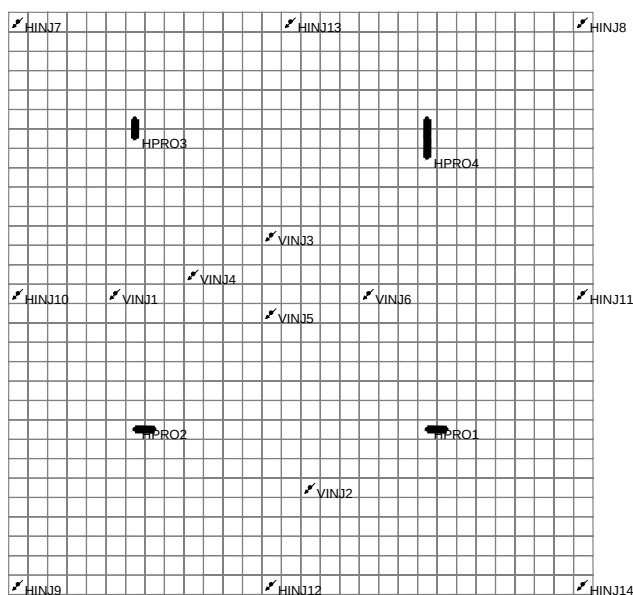


Figura 57. Disposição dos poços após a evolução – Exp 3

Na Figura 57 nota-se que o durante a evolução foi mantida a configuração de poços produtores, apenas incrementando a trajetória do HPRO-4 (produtor localizado acima e à direita). No entanto, a configuração de poços injetores foi modificada incrementando o número de poços injetores e sendo colocados os novos poços em posições próximas ao centro do campo, mantendo uma notória simetria.

O incremento de VPL foi de US\$ 274,646,438.00 o que equivale a um ganho de 19.71%.

5.4.2. Experimentos com reservatório 33x57x3

Nesta seção mostram-se experimentos realizados no reservatório 33x57x3, de características heterogêneas, com emprego de inicialização aleatória, sementes iniciais, e usando informações dadas pelo mapa de qualidade e mapa de aquíferos.

Este modelo de reservatório é mais complexo do que o reservatório 30x30x1, isto se reflete no tempo computacional que cada simulação consome. Por isto, os experimentos realizados o campo heterogêneo foram efetuados utilizando a plataforma de algoritmo genético com avaliações distribuídas dada a necessidade de reduzir o tempo computacional já esperado para as otimizações a serem realizadas..

5.4.2.1. Experimento 1 – Inicialização Aleatória

Este experimento consistiu em evoluir soluções para o reservatório 33x57x3 de características heterogêneas, considerando poços verticais e horizontais e com inicialização aleatória. Os parâmetros usados para o algoritmo genético foram que se mostram na Tabela 33 a seguir:

Parâmetros do AG	Valor
Número de gerações	80
Número de indivíduos	160
Número de rodadas	1
<i>Steady-State</i>	0.6
Probabilidade de poço	0.7
Taxa de <i>crossover</i> (adaptativa)	0.65 – 0.08
Taxa de mutação (adaptativa)	0.10 – 0.50

Tabela 33. Parâmetros do AG: reservatório 33x57x3

As restrições utilizadas neste experimento foram as mostradas na Tabela 34 a seguir.

Restrições no AG	Valor(m)
Distância mínima entre poços	400
Máximo comprimento de poço horizontal	1500

Tabela 34. Restrições nas soluções do AG: reservatório 33x57x3

Após a execução deste experimento, as curvas de evolução *offline*, *online* e de melhor elemento obtidas são as mostradas na Figura 58 a seguir.

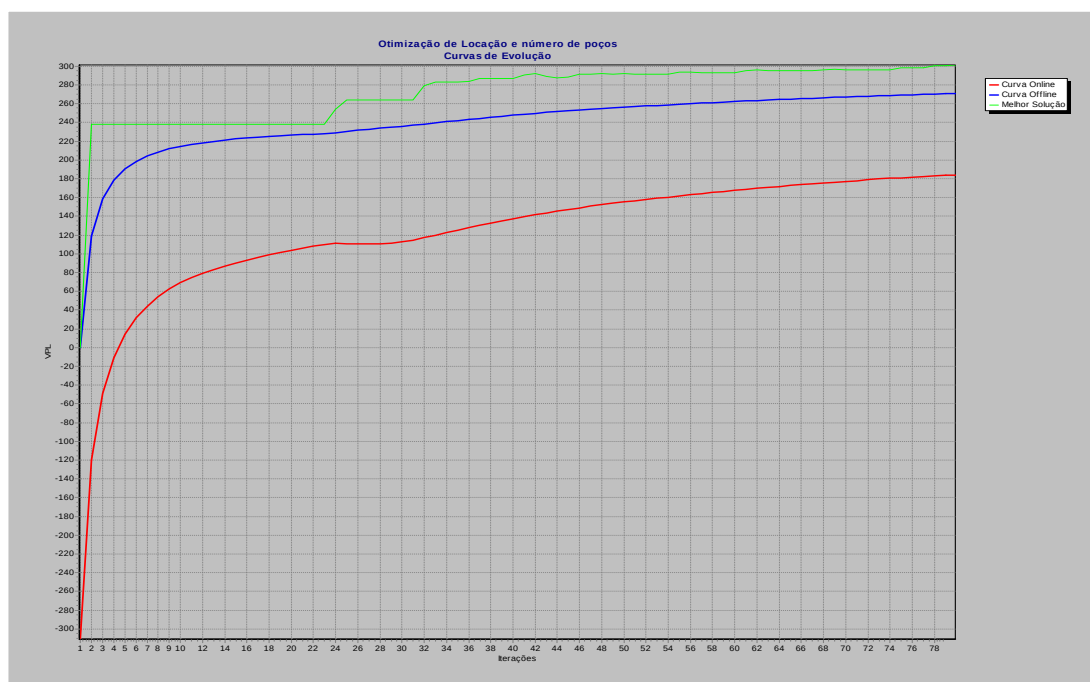


Figura 58. Curvas de evolução: reservatório 33x57x3 – Exp 1.

A melhor solução obtida neste experimento possui os seguintes valores.

VPL **325 988 573,50 (US\$)**
Óleo recuperado **1,413x10⁸ (bbl)**

Na Figura 59, a seguir, mostra-se a disposição de poços obtida pela evolução e, na Figura 60, mostra-se a curva de óleo acumulado obtida pela simulação desta alternativa.

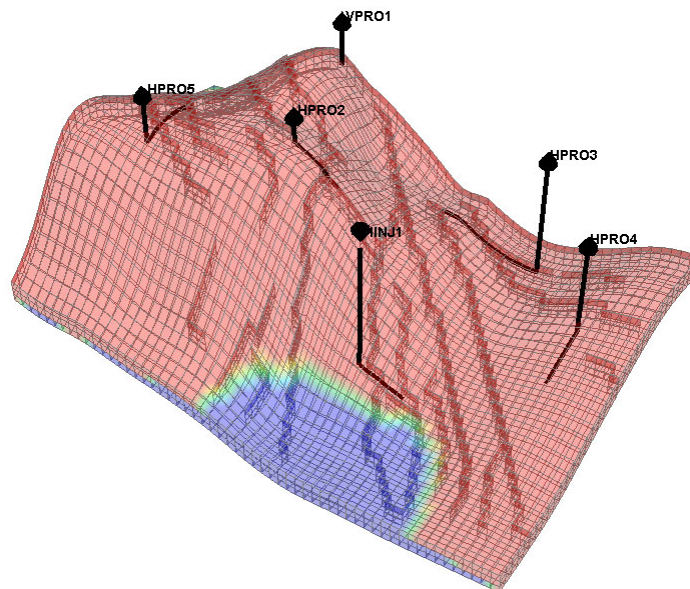


Figura 59. Disposição de poços: reservatório 33x57x3 – Exp 1.

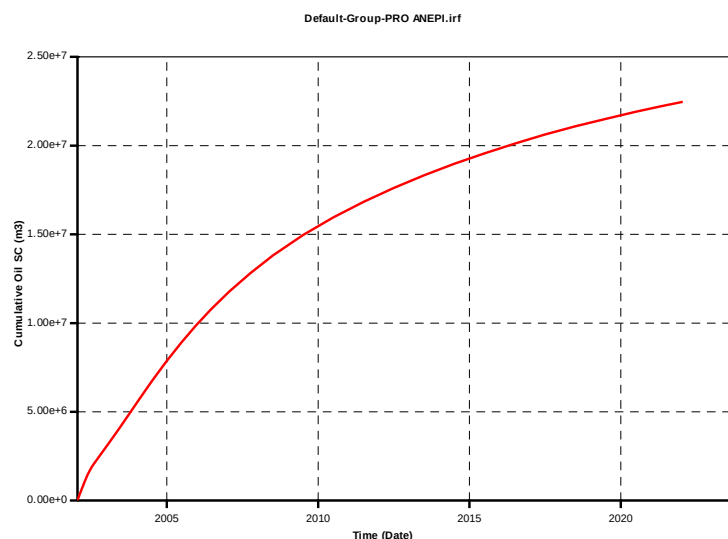


Figura 60. Curva de óleo acumulado: reservatório 33x57x3 – Exp 1.

Observa-se que os poços estão colocados de forma bastante regular respeitando a restrição de distância mínima imposta pelo sistema. Como característica que cabe ressaltar, o sistema otimizador colocou um poço injetor horizontal próximo do aquífero. A curva de produção de óleo acumulado permite visualizar que, até o final da simulação de 20 anos, existe recuperação de óleo.

Este experimento foi realizado utilizando avaliações distribuído em um ambiente paralelo que constou de 42 computadores, como resume a Tabela 35 a seguir

Computadores	Processador
LABORAT. DEE 434 (20)	AMD-XP2600
LABORAT. DEE 430 (19)	AMD-XP1500
ICA SIQURI (01)	AMD-XP2400
ICA LAMBADA (01)	AMD-XP2800
ICA HUAYLAS (01)	AMD-XP2400

Tabela 35. Computadores utilizados: reservatório 33x57x3 – Exp 1

Cada simulação leva em média 35 segundos para ser realizada em um computador com processador AMD-XP2400. Considerando que o número de avaliações realizadas é de 12800, o tempo esperado de otimização em um único computador é de aproximadamente 124 horas (pouco mais de 5 dias). Com os 42 computadores, este tempo diminuiu para 3.66 horas, o que representa uma redução de 34 vezes. Na Figura 61 se mostra um diagrama de barras indicando a quantidade relativa de avaliações que cada computador realizou

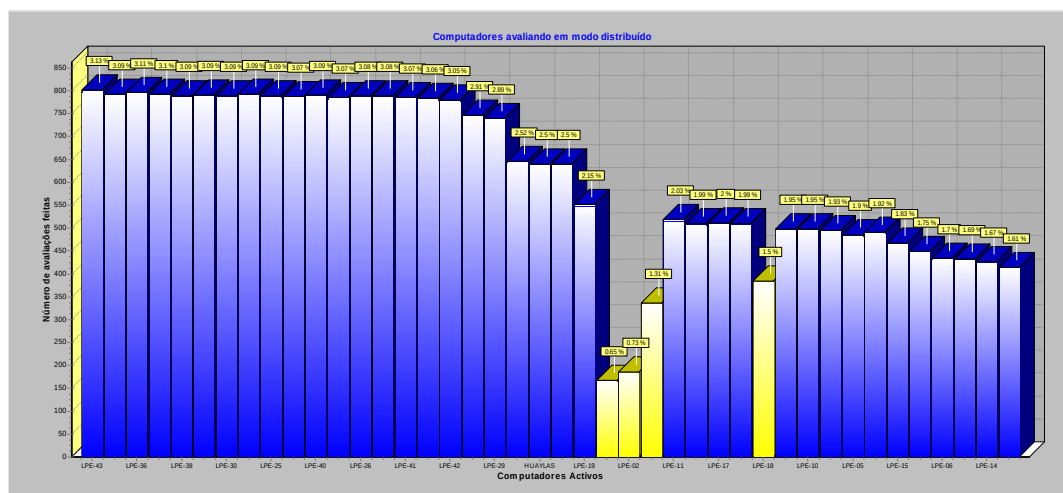


Figura 61. Avaliações por computador: reservatório 33x57x3 – Exp 1.

Neste gráfico, a altura da barra significa o grau de colaboração do micro no conjunto total de avaliações. Nota-se as 17 primeiras barras. Com alturas similares, que correspondem a computadores AMDXP-2600 do laboratório DEE-434, as outras barras, de menor altura, são dos computadores da sala DEE-430 e do laboratório ICA, com processadores de menor poder computacional. Barras amarelas indicam computadores em que ocorreu erro de tipo *time-out* (vide seção 4.8.3 para mais detalhes) durante o processo de otimização e consequentemente, não continuaram a realizar avaliações até o final da otimização, por isso o grau de colaboração destes ficou menor ao final do processo de otimização. Contudo, o sistema de avaliações distribuídas, reenviou

para outros computadores os cromossomas não avaliados corretamente, onde foram efetivamente avaliados.

5.4.2.2.

Experimento 2 – Inicialização Aleatória

Este experimento é uma nova realização com parâmetros de AG e de cálculo do VPL iguais ao Experimento 1, a única diferença existente são os indivíduos criados na inicialização da população, que são diferentes aos do experimento anterior dada a característica aleatória do processo de inicialização da população. O intuito deste experimento é verificar se o AG converge para soluções similares em diferentes execuções. Após a execução do algoritmo, os resultados foram os seguintes.

Melhor cromossoma com

VPL **340 840 442,80 (US\$)**

óleo recuperado **$1,409 \times 10^8$ (bbl)**

na Figura 62 se mostra a alternativa no campo e, na Figura 63, a curva de óleo acumulado.

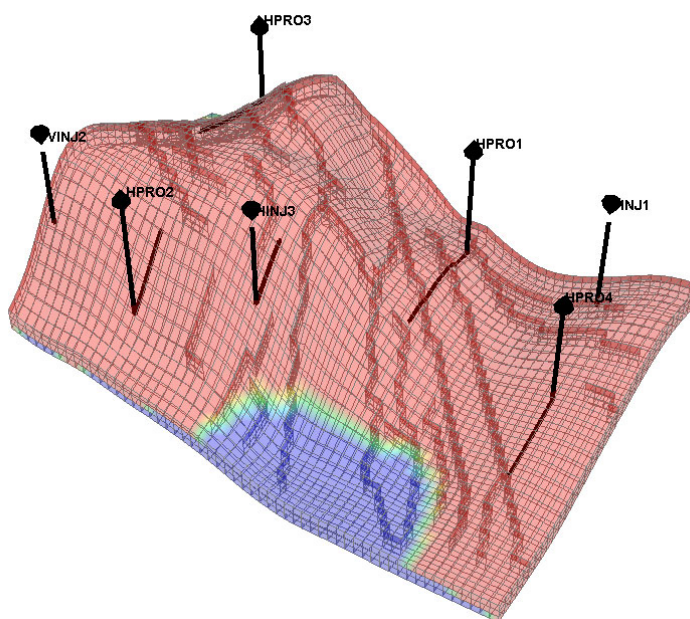


Figura 62. Disposição dos poços: reservatório 33x57x3 – Exp 2.

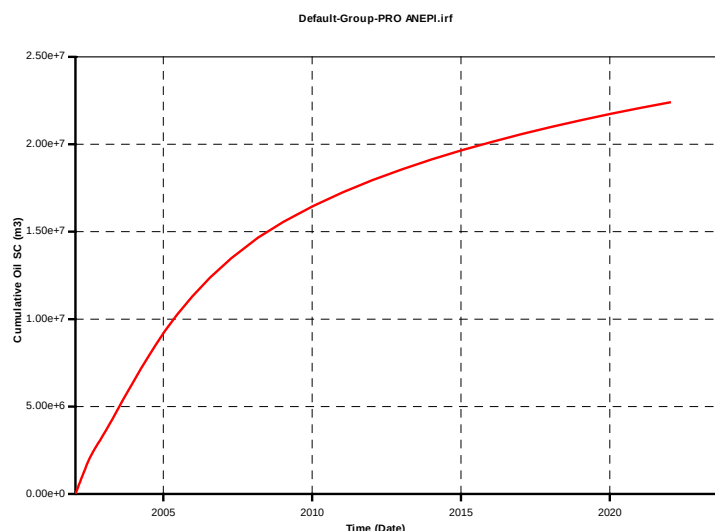


Figura 63. Curva de óleo acumulado: reservatório 33x57x3 – Exp 2.

Neste caso, não houve poço injetor perto do aquífero, mas os poços começaram a se alternar entre injetores e produtores com direções perpendiculares à direção mais longa do campo.

A quantidade de óleo recuperada após os 20 anos de exploração resultou levemente menor, porém, o VPL desta alternativa é maior em US\$14.851.869,30, ou 4,36%. Isto ocorreu porque a alternativa obtida neste experimento recupera mais óleo em tempos iniciais. Pode se observar na curva de óleo recuperado da Figura 63 que, até o ano 2015 a alternativa do experimento 2 recuperou mais óleo. Não entanto, entre 2015 e 2022, a recuperação diminuiu bastante, contudo, o cálculo do VPL, ao colocar mais ênfase na recuperação de óleo em tempos iniciais, permite obter um resultado maior que o VPL da alternativa encontrada no experimento 1.

5.4.2.3.

Experimento 3 – Semente Inicial

Este experimento consistiu em realizar uma execução do algoritmo genético e o uso de sementes iniciais. Estas sementes foram as alternativas obtidas nos experimentos 1 e 2. Os parâmetros do AG, foram alterados aumentando o número de gerações para 200 e diminuindo o número de indivíduos da população para 60. Isto com o objetivo de realizar mais seleções e reproduções sem aumentar muito o número total de avaliações.

As curva de evolução *offline*, *online* e de melhor elemento obtidas neste experimento se mostram na Figura 64 a seguir.

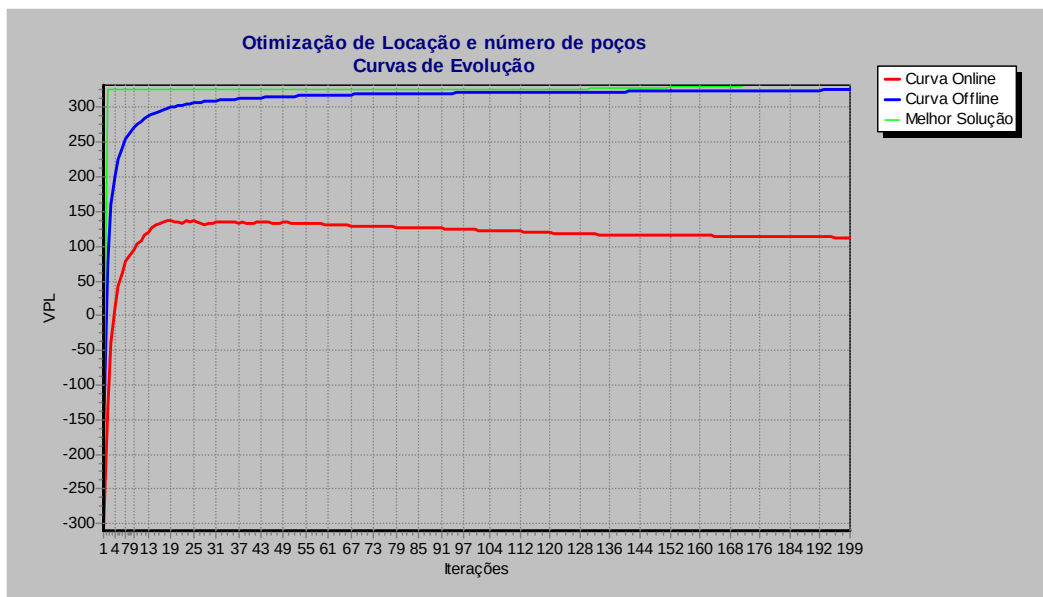


Figura 64. Curvas de Evolução: reservatório 33x57x3 – Exp-03

Após realizadas as 200 gerações, os valores de avaliação da melhor solução obtida foram os seguintes:

VPL **346 633 395,80 (US\$)**
 óleo recuperado **$1,418 \times 10^8$ (bbl)**

A Figura 65 mostra a disposição de poços no campo e, a Figura 66, a curva de óleo acumulado obtida para esta alternativa.

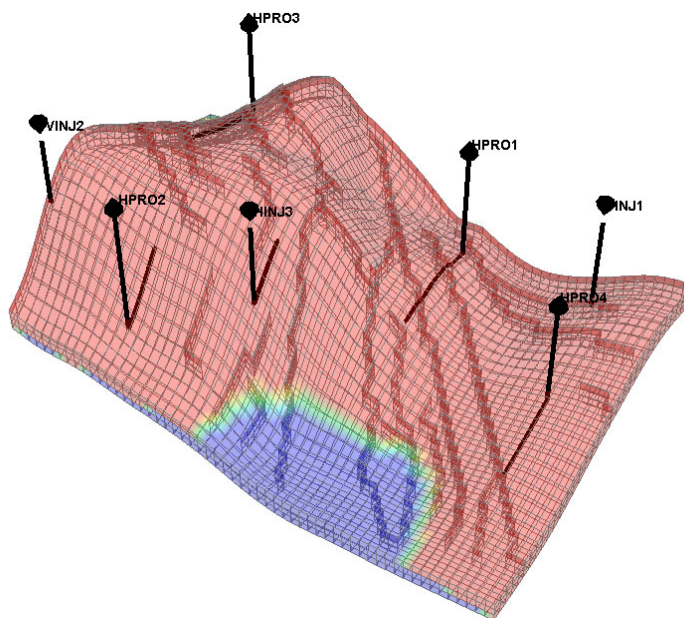


Figura 65. Disposição dos poços: reservatório 33x57x3 – Exp 3.

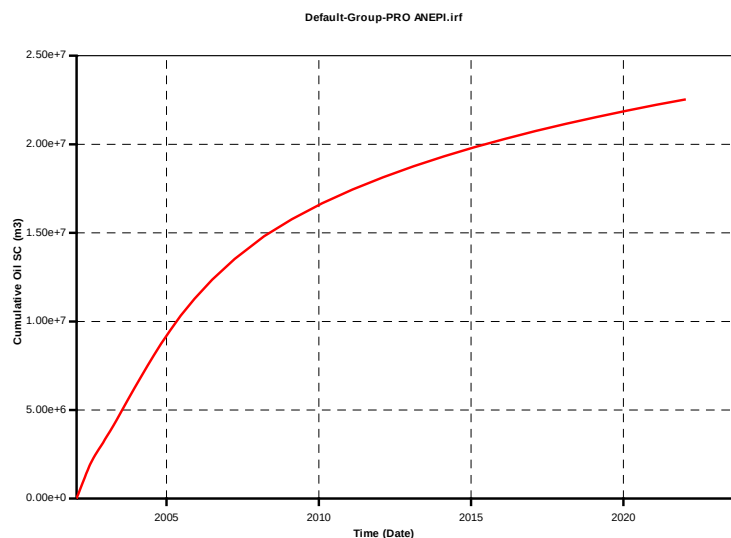


Figura 66. Curva de óleo acumulado: reservatório 33x57x3 – Exp 3.

Nesta alternativa, a disposição de poços ficou similar com a disposição da melhor alternativa do experimento 2, as diferenças se encontram em alguns poços que foram deslocados como se mostra a seguir.

VINJ2	: (7,2)	deslocado para	(8,1)
HPRO4	: (6,12,2)	deslocado para	(5,12,2)
HPRO5	: (30,10,1)	deslocado para	(30, 9, 2)

Nenhum dos poços horizontais foi alterado em comprimento ou direção.

Com estas poucas alterações, o óleo total recuperado nos 20 anos, resultou maior que o valor de óleo acumulado obtido nos experimentos 1 e 2. Assim, a curva de óleo acumulado manteve a forma da curva do experimento 2 com maior recuperação de óleo em tempos iniciais e também aumentou a recuperação de óleo total. O VPL teve um incremento de US\$5,8 milhões, que representa um ganho de 1.7% se comparado com o resultado da experiência 2.

5.4.2.4.

Experimento 4 – Mapas de Qualidade e Aqüíferos

Neste experimento o algoritmo genético empregou a inicialização de indivíduos da população que aproveita a informação dada pelos mapa de qualidade e aqüífero como explicado na seção (4.4).

Os mapas de qualidade e aqüífero foram obtidos seguindo a metodologia descrita em (Cruz 1999, 2000) para uma única realização. A seguir mostram-se os mapas obtidos para este campo:

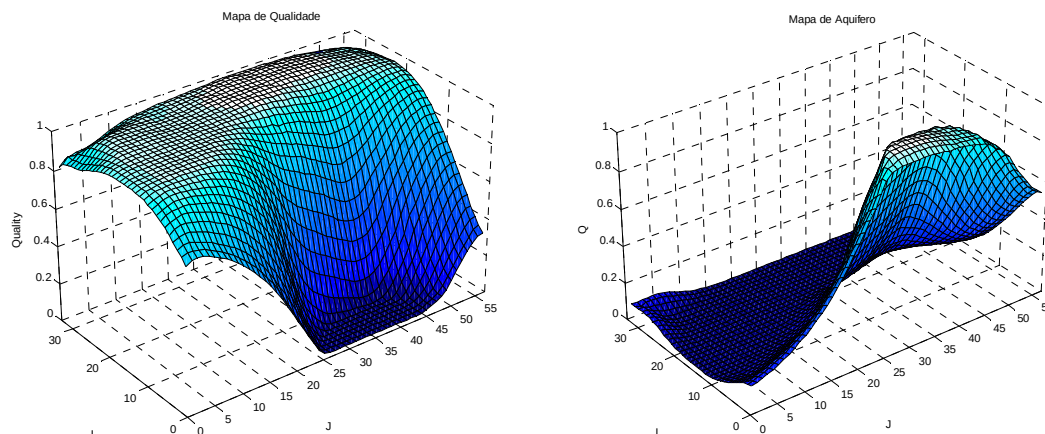


Figura 67. Mapa de qualidade e mapa de aquífero – Exp 4.

Como explicado na seção (4.4), a inicialização que emprega os mapas deve dar grande preferência a colocar poços produtores nas áreas com maior qualidade de óleo e a colocar poços injetores nas áreas de aquífero. Isto é feito para alguns elementos da população inicial do algoritmo genético, o restante é inicializado de forma aleatória.

O melhor indivíduo encontrado teve os seguintes valores de avaliação

VPL **391 387 243,00 (US\$)**

óleo recuperado **$1,617 \times 10^8$ (bbl)**

Na Figura 68 se mostra a disposição de poços no campo e, na Figura 69, a curva de óleo acumulado obtida ao simular esta alternativa.

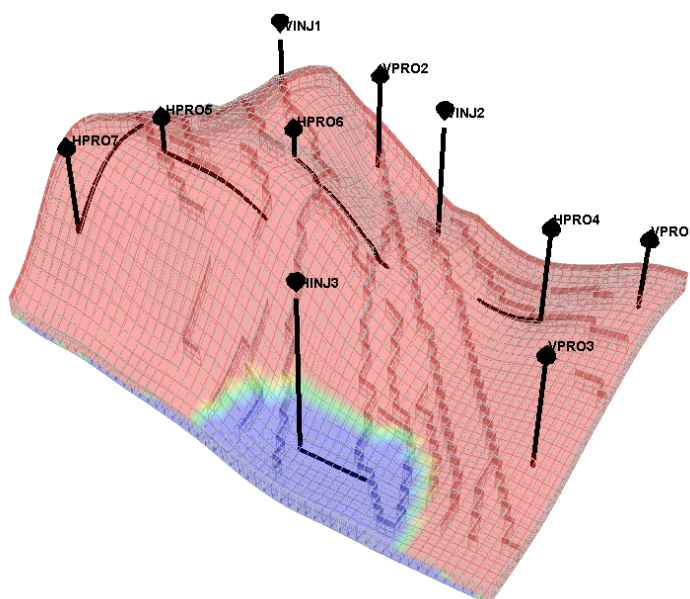


Figura 68. Disposição dos poços: reservatório 33x57x3 – Exp 4.

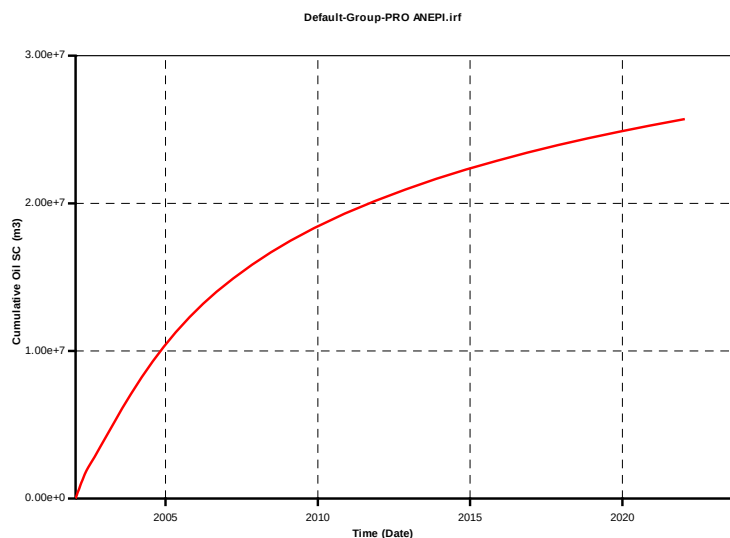


Figura 69. Curva de óleo acumulado: reservatório 33x57x3 – Exp 4.

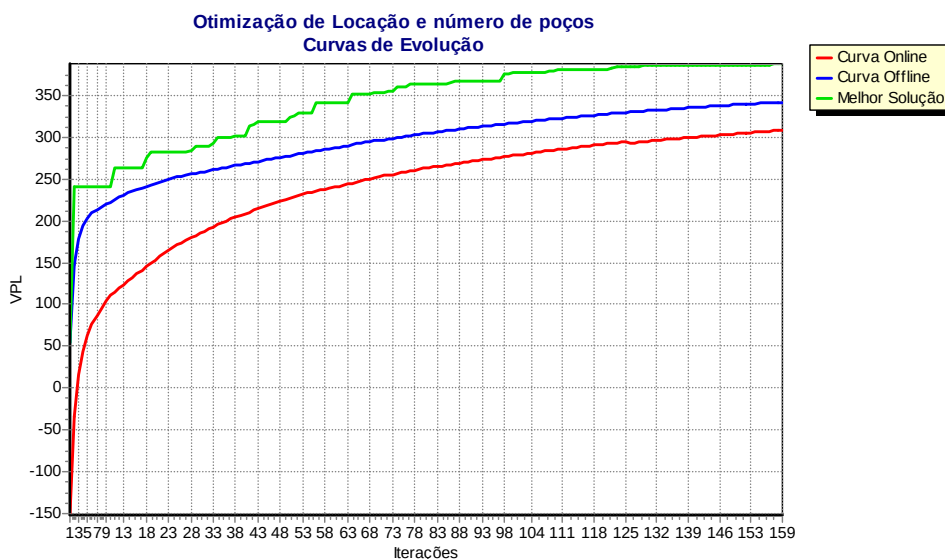


Figura 70. Curvas de evolução: reservatório 33x57x3 – Exp 4.

A solução encontrada neste experimento é claramente superior às obtidas nos experimentos anteriores. O valor de VPL é superior ao VPL da experiência 3 em US\$44.753.847,20. Isto significa um ganho de 11.43%.

Uma característica que cabe ressaltar é o fato do algoritmo genético ter mantido o poço injetor horizontal no setor de aquífero que foi colocado segundo o critério dado pelo mapa de aquífero.

5.4.3. Uso do módulo de inferência na otimização

Nesta seção, mostram-se os experimentos realizados no sistema de otimização utilizando os modelos de inferência das curvas de produção de óleo.

Para realizar estes experimentos foram levadas em conta as seguintes condições

- Limitar as evoluções do sistema otimizador para empregar apenas poços verticais e permitir como máximo 10 poços, dado que os modelos de inferência desenvolvidos suportam atualmente até 10 poços verticais;
- A cada final da geração deve ser avaliado o melhor indivíduo, utilizando sempre o simulador de reservatório;
- Utilização do campo heterogêneo, dado que os aprendizados foram realizados com este campo.
- Taxa de uso de simulador $\gamma = 0.7$
- Emprego da arquitetura de avaliações distribuídas com 7 computadores do laboratório ICA, como se descreve na tabela a seguir

Computadores	Processador
ICA BOLERO (01)	AMD-XP2800
ICA POLCA (01)	AMD-XP2400
ICA SIQURI (01)	AMD-XP2400
ICA LAMBADA (01)	AMD-XP2800
ICA DIABLADA (01)	INTEL-P4HT 3.0
ICA SAMBA (01)	INTEL-P4 2.6G

Tabela 36. Computadores utilizados: módulo de inferência da produção

5.4.3.1. Experimento 1 – Pontos de Curva de Produção

Neste experimento foi usado o modelo de aproximação de pontos de curva de produção descrito na seção (4.5.2.1) utilizando o modelo *neuro-fuzzy* hierárquico (NFHB), o algoritmo genético foi estabelecido para realizar 100 gerações, com 200 indivíduos na população, e as mesmas taxas de operadores de cruzamento e mutação empregadas nos experimentos de otimização.

Após a realização deste experimento, curvas de evolução foram obtidas e são mostradas na figura a seguir.

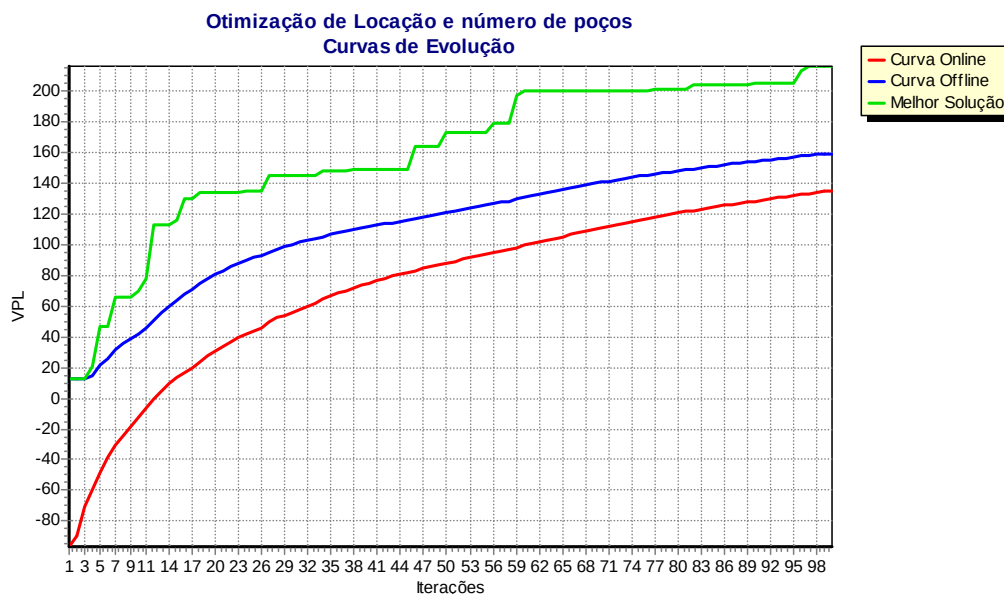


Figura 71. Curvas de evolução: módulo de inferência da produção – Exp 1.

Após as 100 gerações, o valores de avaliação da melhor solução obtida foram os seguintes.

VPL	216 176 554,10 (US\$)
Óleo recuperado	1,265x10⁸ (bbl)
Tempo de execução	14:48 horas
Tempo de execução sem inferência	19.15 horas
Redução de tempo (%)	25.00%

O tempo de execução resultou menor ao tempo que levaria o mesmo processo de otimização sem o uso da inferência, isto se deve ao uso médio do simulador em 0.7 das avaliações em total. A redução de tempo esperada é 30%, porém, a redução encontrada foi de 25.00%.

Esta perda de redução pode ser explicada pelos seguintes fatos

- o processo de recuperação de respostas (*recall*) dos módulos de inferências leva um tempo para ser efetuado;
- existe ainda o tempo que leva o cálculo do VPL;
- ao finalizar a avaliação de uma dada geração, o melhor elemento encontrado é reavaliado usando o simulador de reservatórios, isto antes de aplicar o *steady-state* ou elitismo.

O valor VPL e óleo recuperado mostram-se menores aos obtidos nos experimentos de otimização da seção (5.4.2). Isto ocorre, porque os experimentos foram realizados com condições impostas para poder utilizar os

módulos de inferência, a mais relevante destas condições é a limitação de tipo de poço para ser apenas poços verticais.

5.4.4.

Eficiência do uso de inferência de curvas de produção

Este experimento foi realizado para verificar a relação entre a taxa de uso do simulador e a redução de tempo obtida. O teste consistiu em executar um algoritmo genético distribuído com 100 gerações e 200 indivíduos na população sob o ambiente de 7 computadores da rede ICA mostrados na Tabela 36 para avaliações distribuídas. Neste experimento variou-se a taxa de uso de simulador decrescentemente desde 1.0 (uso total do simulador) até o valor de 0.0 (uso total do módulo de inferência). Na figura se mostra a relação encontrada.

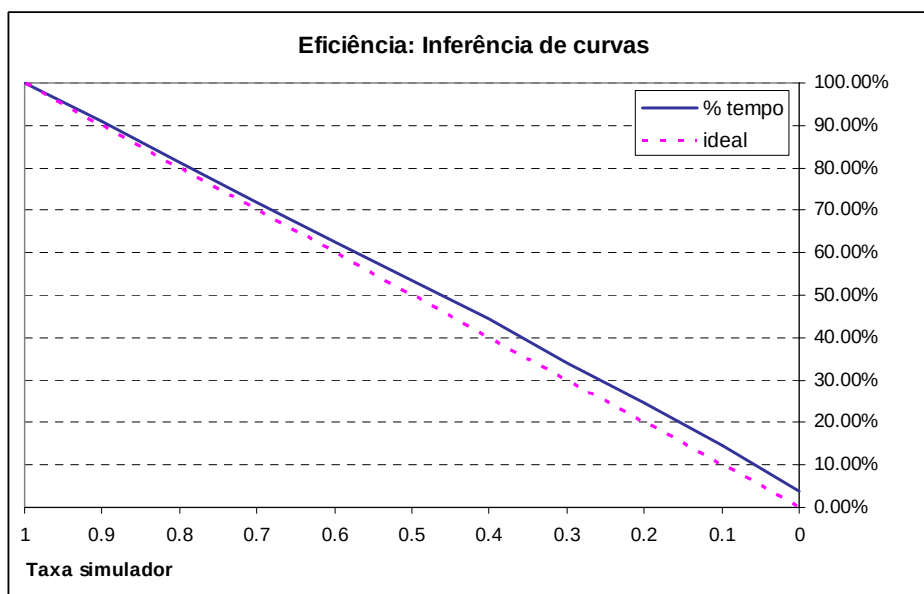


Figura 72. Eficiência do uso de inferência de curvas.

A Figura 72 mostra como o sistema reduz o tempo necessário para a execução da otimização especificada acima. A medida que a taxa de uso de simulador diminui, o sistema utiliza mais vezes o módulo de inferência. Para um caso de taxa de uso de simulador $\gamma = 0.0$, o sistema emprega apenas o módulo de inferência. Pode-se notar neste caso que o tempo ficou reduzido até aproximadamente um 4% do valor inicial, o que significa um ganho de tempo de 25 vezes. Todavia, deve-se levar em consideração que o ganho de tempo obtido foi para uma configuração que já teve um ganho de tempo dado pelo uso de avaliações distribuídas, o que leva a ter um ganho total de 135 vezes se comparado com um processo iterativo em um único computador.

5.4.5. Eficiência do uso de computação distribuída

Este experimento foi realizado com o objetivo de verificar a relação entre o número de computadores colocados no ambiente paralelo e o ganho de tempo obtido comparado com o ganho esperado para uma rede ideal. O teste consistiu em executar repetidamente um algoritmo de 20 gerações e 100 indivíduos na população, empregando desde um único computador até 42 computadores das redes dos laboratórios ICA e DEE. Estes laboratórios constam na época de 19 computadores AthlonXP2600 da rede DEE1, 19 computadores AthlonXP1500 da rede DEE2 e 4 computadores AthlonXP2200, 2400, 2400 e 2800 da rede ICA. Na Figura 73, se mostra um gráfico com a relação encontrada.

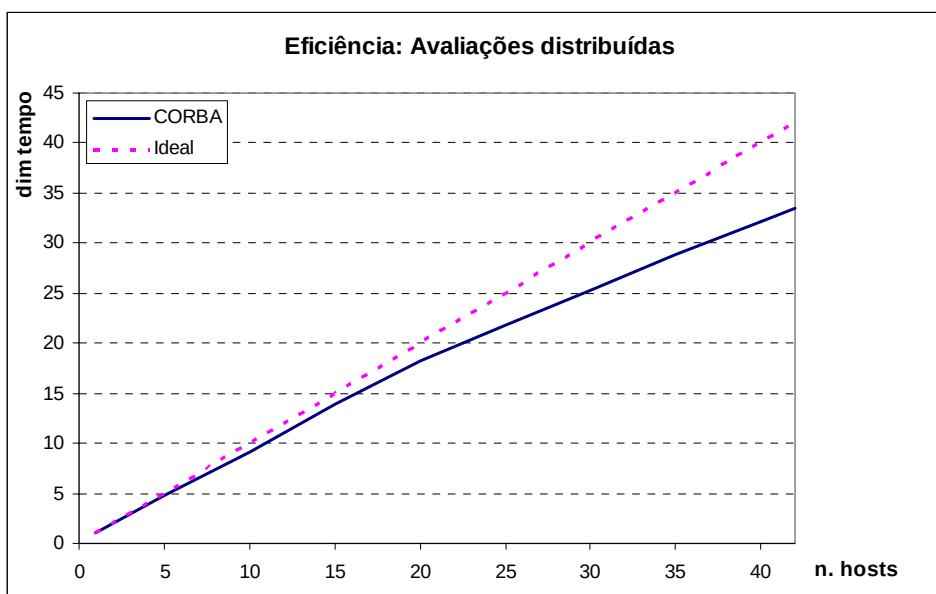


Figura 73. Eficiência do uso da computação distribuída

Na Figura 73 pode se ver que, na medida que aumenta o número de computadores no ambiente paralelo, a diferença das reduções de tempo ideais e encontrada também aumenta, isto se deve aos seguintes fatos

- maior custo de comunicação existente;
- custo computacional adicional que leva o gerenciamento de maior número de *threads* no processo *master*, onerado no sistema operacional;
- diferenças de velocidade existentes entre os computadores existentes nas redes ICA, DEE1 e DEE2, o conjunto total da computadores utilizados não é perfeitamente homogêneo

5.5. Discussão

Dos resultados obtidos nos experimentos realizados neste trabalho, observa-se que, os erros apresentados nos modelos de aproximação de funções são baixos. Para alternativas contendo poucos poços, o que significa poucas variáveis de entrada para os modelos de aproximação, o melhor desempenho foi obtido pelo modelo *Neuro-Fuzzy* Hierárquico, que alcançou valores de erro MAPE menores do que 1% tendo um desempenho bastante melhor do que os modelos de redes neurais. Para alternativas contendo mais poços, isto é, um maior número de variáveis de entrada, os valores de erro ficaram menores de 5% para todos os casos. Estes resultados podem ser vistos desde a Tabela 12 até a Tabela 22. Os sistemas de aproximação, quando treinados, servem como *proxy* de simulação para o módulo otimizador, podendo fornecer resultados bem aproximados de forma mais imediata que utilizando o simulador. Um problema persistente nesta abordagem é o grande investimento inicial necessário para criar os conjuntos de entradas e saídas que acarreta o uso intensivo do simulador de reservatórios.

Durante os experimentos com o sistema otimizador foram encontrados resultados muito interessantes: partindo de inicialização aleatória de indivíduos para a população inicial, foram alcançadas alternativas com valores VPL bastante altos. Resultados ainda mais relevantes no que se refere a VPL foram encontrados ao colocar conhecimento prévio ao algoritmo genético como sementes iniciais. No experimento 3, com campo homogêneo, a semente inicial foi uma configuração clássica (*4-five-spots*) e o algoritmo foi capaz de descobrir a necessidade de evoluir o sistema de injeção colocando vários poços injetores verticais sem alterar a configuração produtora dada na semente inicial. O incremento no VPL resultante alcançou 19.71%; além disso, a configuração de poços resultante deste experimento mostra uma tendência para um *layout* simétrico de poços, isto é esperado, dada a característica homogênea do campo. Contudo, como esperado, as respostas da evolução que partiram de inicialização aleatória tiveram resultados com VPL superior ao das evoluções que partiram de conhecimento inserido em sementes iniciais.

No que diz respeito ao emprego de informação da caracterização por mapas de qualidade, para o campo homogêneo não foram realizados estes experimentos, dado que para este campo, o mapa de qualidade e mapa de

aquífero também possuem características homogêneas, isto é, as qualidades obtidas são todas similares (lembrando uma distribuição de probabilidades uniforme que representa ausência de conhecimento), portanto não existe informação a ser aproveitada na otimização.

No entanto, nos experimentos que empregaram a caracterização realizados com a configuração de reservatório de características heterogêneas, foram encontrados resultados que também alcançam VPLs altos, sendo que, o uso de informação proveniente dos mapas de qualidade e mapas de aquífero permitiu obter resultados com VPL ainda superiores, ultrapassando os valores encontrados a partir de inicialização aleatória pura e os alcançados a partir de sementes iniciais achadas previamente pelo mesmo algoritmo otimizador.

O campo heterogêneo, dada sua maior complexidade, exige maior tempo computacional para realizar a simulação de uma alternativa qualquer. Nos experimentos que envolveram o uso do campo heterogêneo foi benéfico em tempos computacionais ter o ambiente paralelo para distribuir as avaliações realizadas durante o processo de otimização. A redução de tempo de otimização (*speed-up*) foi evidente e mostra que o poder computacional de um conjunto de processadores poder ser bem aproveitado em um sistema de otimização que realize muitas iterações como ocorre com os sistemas baseados em algoritmos evolucionários. Durante a execução de alguns experimentos ocorreram erros durante a avaliação dos cromossomas em algum dos computadores remotos, o sistema foi capaz de detectar estes erros e reenviar o cromossoma para outro computador sem prejudicar muito o desempenho global do processo de otimização.

Os experimentos que envolveram o uso do módulo de aproximação de funções para inferir a curva de produção de óleo, também mostraram que pode ser obtido um grande ganho de tempo adicional ao já obtido ao empregar avaliações distribuídas.