

5

Simulação Monte Carlo

A variedade de metodologias para a estimação de Efeito de Tratamento sob Ignorabilidade, bem como de possíveis implementações de cada uma, motivou, nos últimos anos, uma série de estudos de simulação comparando o desempenho das alternativas em amostras finitas. Neste capítulo, revisaremos a literatura de simulações relacionadas à questão da combinação de métodos proposta ao longo deste trabalho. Em seguida, realizaremos um novo estudo de simulações com o objetivo de avaliar a relevância dos estimadores propostos na seção 4.3.

5.1

Resultados da Literatura de Simulações sobre Combinação de Métodos

Estudos de simulação considerando métodos combinados foram motivados pelos estudos teóricos sobre esses estimadores e, devido a isto, são recentes. Uma parte considerável dessa literatura foi desenvolvida pelos próprios proponentes dos procedimentos duplamente robustos, (Bang e Robins (2005), Robins et al. (2007), Neugebauer e Van der Laan (2002)), como forma de divulgá-los. Os resultados publicados até o presente são nitidamente divergentes entre si.

Lunceford e Davidian (2004) estudam o comportamento em amostras finitas de diversos estimadores baseados em *propensity score*. O foco do trabalho é a comparação entre métodos de Reponderação e de estratificação da amostra, com vantagem para a primeira abordagem, nas simulações apresentadas. Um ponto interessante para nossa discussão é que tanto Reponderação quanto estratificação apresentam melhor desempenho quando ajustadas por uma regressão. A recomendação final é o uso de um estimador duplamente robusto na forma de Reponderação aumentada como proposto por Scharfstein, Robins e Rotnitzky (1999).

Bang e Robins (2005) apresentam extensões da estimação duplamente robusta para diferentes contextos envolvendo dados faltantes. Nas simulações, são destacadas as situações em que os modelos usados para estimar a probabilidade de observação e os valores potenciais estão mal especificados. O principal resultado é o bom desempenho dos estimadores duplamente robustos, mesmo

quando ambas as especificações estão incorretas.

Num trabalho semelhante, Kang e Schaeffer (2007) discutem métodos duplamente robustos e os comparam com métodos baseados em Regressão e Reponderação. Também é dada ênfase o comportamento dos diferentes procedimentos quando as estimações preliminares falham, de forma limitada, em captar os modelos populacionais de seleção e regressão. Os resultados obtidos mostram sensibilidade dos métodos que usam *propensity score* à má especificação dos modelos para este objeto. Diferentemente do que ocorre em Bang e Robins (2005), os melhores desempenho são dos estimadores simples de Regressão, enquanto os duplamente robustos representam melhoria em relação aos de Reponderação em alguns casos. Num comentário sobre este trabalho, Robins et al. (2007) questionam os resultados, argumentando que peculiaridades das simulações favorecem particularmente os estimadores de Regressão. A fragilidade dos resultados de Kang e Schaeffer é sugerida pela consideração do modelo modificado pela inversão do indicador T , que apresenta resultados divergentes em relação aos do modelo original.

Busso, DiNardo e McCrary (2009) comparam estimadores de Reponderação, *Matching* e o estimador duplamente robusto em forma de Regressão Ponderada. Uma característica peculiar dos modelos adotados é que as funções de regressão dependem de X apenas via $p(X)$. Em suas simulações, os estimadores de Reponderação apresentam os melhores resultados, exceto quando a hipótese de Sobreposição é violada, ou em situações próximas, em que o *propensity score* atinge valores próximos de zero ou um. Neste caso, todos os métodos apresentam considerável viés e pouca semelhança com as previsões da teoria assintótica. Os autores ressaltam a importância desta hipótese, e apontam a sua falha como razão de resultados inusitados em outros estudos de simulação.

5.2

Um Novo Estudo de Simulações

O estudo de simulações a seguir tem o objetivo de avaliar a interação entre elementos do modelo populacional e o desempenho da combinação de Reponderação e Imputação em relação ao uso de apenas uma das técnicas, além de suas diferentes formas de implementação. Neste sentido, quatro elementos principais serão analisados.

O primeiro aspecto que buscamos ressaltar é o efeito da forma funcional das funções de regressão e do *propensity score*, particularmente quanto à sua suavidade. Pela discussão do capítulo anterior, a motivação para utilizar o inverso do *propensity score* como peso numa regressão é eliminar a correlação

entre os regressores e variáveis omitidas. A má especificação pode ser incluída neste contexto, se interpretada como a omissão de uma parte de $m_t(X)$. Portanto, funções de regressão pouco suaves, que projetadas numa pequena base das variáveis auxiliares deixam um resíduo considerável, devem implicar ganhos importantes em reponderar previamente a amostra.

Em segundo lugar, são estudados os efeitos da heterocedasticidade sobre a conveniência de realizar Regressões Ponderadas. Uma vez que a reponderação da regressão em $\hat{\beta}_{RP}$ segue um critério completamente independente da eficiência dessa estimação, o efeito sobre a precisão do estimador pode ser adverso.

Terceiro, será considerado o efeito da quantidade de dimensões do espaço das variáveis auxiliares. A discussão do artigo de Robins e Ritov (1997) sugere que a ‘maldição da dimensionalidade’ tem mais chances de comprometer o desempenho de estimadores não robustos. Por outro lado, vimos que a hipótese de Ignorabilidade exige do econometrista um esforço em obter toda a informação possível para controlar a heterogeneidade, o que deve refletir no número de variáveis auxiliares. Apesar disto, a maioria das simulações disponíveis na literatura emprega X univariada. Dos estudos mencionados anteriormente, apenas o de Lunceford e Davidian (2004) expressam esta preocupação.

Finalmente, avaliaremos o uso do *propensity score* verdadeiro como elemento do cálculo dos pesos na Regressão Ponderada e do termo de correção na Reponderação Aumentada. Conforme apontamos no capítulo anterior, é possível que um estimador Duplamente Robusto que empregue o valor populacional de $p(\cdot)$, no lugar de uma estimativa, atinja o Limite de Eficiência Semiparamétrico. Isto indicaria que a incorporação do conhecimento disponível sobre essa função pode ser feita de maneira simples, reduzindo a dificuldade computacional do procedimento sem prejuízo à variância assintótica.

5.3

Estimadores e Modelos

Consideraremos estimadores de Regressão ($\hat{\beta}_{rep}$), Imputação ($\hat{\beta}_{imp}$) e duplamente robustos dos tipos $\hat{\beta}_{FI}$ (Reponderação Aumentada ou Função Influência) e $\hat{\beta}_{RP}$ (Regressão Ponderada). As estimativas preliminares de $p(\cdot)$ serão realizadas por *sieve*/logit, para todos os estimadores. As funções de regressão $m_t(\cdot)$, $t = 0, 1$ serão estimadas por *sieve*/mínimos quadrados ordinários, para Regressão e $\hat{\beta}_{FI}$, e por *sieve*/mínimos quadrados ponderados (com os pesos descritos na seção 4.3), para o uso em $\hat{\beta}_{RP}$. Além disso, calculamos também os estimadores $\tilde{\beta}_{FI}$ e $\tilde{\beta}_{RP}$, análogos a $\hat{\beta}_{FI}$ e $\hat{\beta}_{RP}$, porém

combinando o *propensity score* verdadeiro e as funções de regressão estimadas.

Como forma de estudar a seleção do número de termos das bases dos estimadores *sieve*, consideraremos todo o intervalo de valores $L = 1, 2, 3, \dots, \bar{L}$ para a estimação de $m_t(\cdot)$ e $K = 1, 2, 3, \dots, \bar{K}$ para a estimação de $p(\cdot)$. As bases consistem de polinômios das variáveis auxiliares, em ordem (i) crescente no grau, (ii) decrescente no maior expoente, e (iii) decrescente no expoente da variável de menor número de ordem. Como observado no capítulo anterior, os estimadores de Regressão e Imputação são casos particulares dos estimadores duplamente robustos, respectivamente, quando *propensity score* e funções de regressão são estimados por uma constante. Logo o cálculo $\hat{\beta}_{FI}$ e $\hat{\beta}_{RP}$ inclui o de $\hat{\beta}_{rep}$, quando $L = 1$ e $\hat{\beta}_{imp}$, quando $K = 1$. Além disso, para $L = 1$, $\tilde{\beta}_{FI}$ e $\tilde{\beta}_{RP}$ são iguais ao estimador (ineficiente) de Reponderação pelo *propensity score* populacional.

Para destacar o efeito da suavidade de $p(\cdot)$ e $m_t(\cdot)$, consideramos inicialmente um modelo com uma única variável auxiliar X_i , distribuída uniformemente entre -1 e 1 , e

$$\begin{aligned} Y_i(1) &= m_1(X_i) + u_{1,i} \\ Y_i(0) &= m_0(X_i) + u_{0,i} \\ Y_i &= T_i(m_1(X_i) + u_{1,i}) + (1 - T_i)(m_0(X_i) + u_{0,i}) \end{aligned}$$

com Y_i homocedástica, ou seja, $u_i = (u_{0,i}, u_{1,i})$ independente e identicamente distribuída conforme uma distribuição normal bivariada de média zero e variância σ^2 . A variação das especificações de $p(\cdot)$ e $m_t(\cdot)$ é feita pela manipulação dos coeficientes das seguintes representações:

$$\begin{aligned} \log \left(\frac{p(x)}{1 - p(x)} \right) &= \sum_{k=0}^{\infty} a_k \lambda_k(x) \\ m_1(x) - m_0(x) &= \sum_{k=0}^{\infty} b_k \lambda_k(x) \\ m_0(x) &= \sum_{k=0}^{\infty} c_k \lambda_k(x), \end{aligned}$$

onde λ_k é o polinômio de Legendre de grau k . Os polinômios de Legendre são construídos de forma a serem ortogonais entre si, em relação ao produto interno $\langle f, g \rangle = \int f g dx$; portanto definem funções não correlacionadas entre si de uma variável com distribuição uniforme. O primeiro polinômio é constante ($\lambda_0(x) \equiv 1$), logo, por ortogonalidade, os demais têm média zero, e portanto o

Efeito Médio de Tratamento é dado por

$$\beta = E[Y_i(1) - Y_i(0)] = \sum_{k=0}^{\infty} b_k E[\lambda_k(x)] = b_0.$$

A velocidade de decaimento dos coeficientes das séries simula a suavidade das funções correspondentes. De fato, o papel das hipóteses de suavidade na teoria assintótica dos estimadores estudados é assegurar determinadas taxas de convergência do estimador não paramétrico do primeiro passo. No caso da estimação por *sieve*, a consequência relevante da suavidade é garantir a existência de uma sequência de aproximações (não aleatórias) $\sum_{k=0}^K \tilde{a}_k^K \lambda_k(x)$ convergindo a determinada taxa (em termos de K) para a função estimada (Newey, 1997). Este é justamente o efeito simulado através da manipulação de especificações realizada nos modelos 1 a 3.

No modelo 1, utilizamos $\sigma^2 = 9,5114I_2$ (I_2 é a matriz identidade 2 por 2); $a_0 = 0,5$; $a_k = \frac{k^{-1}}{2}$, se $1 \leq k \leq 10$; $a_k = 0$, se $k > 10$; $b_0 = 1$; $b_k = 20a_k$, se $k \geq 1$; $c_k \equiv 1$. Com esta escolha de parâmetros, o *propensity score* é uma transformação monotônica da função de regressão, o que intuitivamente parece tornar redundantes Reponderação e Imputação, e inócua a combinação entre as duas. O modelo é homocedástico, o que tornaria atrativo o estimador de regressão corretamente especificado.

Nos modelos 2 e 3, são alteradas as especificações de modo que os coeficientes se distribuem diferentemente. No primeiro caso, o decaimento é mais rápido: $a_k = 1,0922 \frac{k^{-2}}{2}$, se $1 \leq k \leq 10$; no segundo, não há decaimento: $a_k = \frac{0,5926}{2}$, se $1 \leq k \leq 10$. É feito o correspondente ajuste para manter $b_k = 20a_k$, se $k \geq 1$, em cada caso. As constantes multiplicadas na redefinição de a_k mantêm o limite de eficiência semiparamétrico igual ao do modelo 1.

Os modelos 4 e 5 introduzem heterocedasticidade. São consideradas as especificações do modelo 1 para a_k , b_k e c_k , mas u_i são independentemente distribuídas com média zero e variância $\sigma^2(X_i)$. No modelo 4, definimos $\sigma^2(X_i) = 20,7767 \text{diag}(1-p(X_i), p(X_i))$ ($\text{diag}(d_1, d_2)$ é a matriz que tem d_1 e d_2 na diagonal principal e zero nas demais entradas), uma especificação favorável à Regressão Ponderada, pois, por exemplo, observações tratadas com *propensity score* alto, que receberão pesos baixos, têm variância maior. Se os pesos são corretamente especificados, correspondem àqueles que otimizam as regressões de $Y(t)$ em X por mínimos quadrados ponderados. No modelo 5, a situação é oposta, com $\sigma^2(X_i) = 17,5421 \text{diag}(p(X_i), 1-p(X_i))$. As constantes nas expressões para $\sigma^2(X_i)$ são introduzidas para preservar o mesmo limite de eficiência semiparamétrico, como fizemos nos modelos 2 e 3.

Para a análise do efeito da dimensão de X_i , o modelo 6 considera três

componentes para esta variável, todas identicamente distribuídas, uniformemente no intervalo $[-1, 1]$, e independentes entre si. Os modelos de seleção para o tratamento e de média condicional são descritos abaixo:

$$p(x) = 0,8x_3L(x_1) + (1 - 0,8x_3)L(x_2)$$

$$L(\cdot) = \frac{\exp(\cdot)}{1 - \exp(\cdot)}$$

$$m_1(x) - m_0(x) = 5 \exp(2x_3)(1.5 + \sin(\frac{\pi}{4}(x_1 + x_2)) - (x_1 - x_2)^2)$$

$$m_0(x) \equiv 0.$$

O par de valores potenciais da variável de interesse tem o componente aleatório u_i independente e identicamente distribuído de acordo com uma distribuição normal bivariada de média zero e variância $10I_2$.

5.4

Resultados

Para as simulações com X unidimensional, foram consideradas amostras com $N = 100$ e $N = 200$ observações. Acima destes valores, os estimadores apresentaram desempenhos indistinguíveis. As tabelas 1 a 4 reportam os erros quadráticos médios simulados dos estimadores, para as duas formas de combinar métodos, com a especificação do modelo 1.

Tabela 5.1: Modelo 1, Regressão Ponderada, N=100

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	3,578	1,004	0,906	0,869	0,914	1,015	1,969	3,767	> 10
2	0,957	1,047	0,900	0,864	0,899	1,003	1,952	4,420	> 10
3	0,890	0,880	0,906	0,867	0,902	0,998	1,961	4,541	> 10
4	0,889	0,862	0,868	0,870	0,902	1,006	1,914	4,520	> 10
5	0,900	0,855	0,866	0,857	0,905	1,012	1,922	4,219	> 10
6	0,911	0,855	0,865	0,855	0,901	1,012	1,952	4,745	> 10
7	0,930	0,863	0,873	0,860	0,906	1,015	1,928	4,670	> 10
8	0,949	0,868	0,882	0,865	0,908	1,015	1,984	4,639	> 10
9	0,991	0,882	0,899	0,876	0,923	1,023	2,046	4,624	> 10
10	1,018	0,930	9,286	> 10	> 10	> 10	> 10	> 10	> 10
11	1,070	0,917	0,941	3,951	> 10	> 10	> 10	> 10	> 10
12	1,120	1,114	> 10	> 10	> 10	> 10	> 10	> 10	> 10

N=100, Limite de Eficiência Semiparamétrico = 0,831, 5000 replicações

É notável que o melhor estimador, em todos os casos, utiliza mais de um termo tanto para estimar o *propensity score* ($K > 1$) quanto as funções de regressão ($L > 1$). Contrastado com a observação na seção anterior, de

Tabela 5.2: Modelo 1, Reponderação Aumentada, N=100

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	3,578	1,004	0,906	0,869	0,914	1,015	1,969	3,767	> 10
2	0,957	1,062	0,902	0,869	0,914	1,015	1,969	3,767	> 10
3	0,890	0,881	0,912	0,871	0,913	1,015	1,969	3,767	> 10
4	0,889	0,867	0,872	0,875	0,915	1,016	1,969	3,767	> 10
5	0,900	0,860	0,874	0,859	0,918	1,018	1,969	3,767	> 10
6	0,911	0,859	0,873	0,858	0,911	1,017	1,968	3,767	> 10
7	0,930	0,867	0,880	0,864	0,917	1,017	1,965	3,766	> 10
8	0,949	0,873	0,888	0,870	0,923	1,020	1,968	3,767	> 10
9	0,991	0,887	0,905	0,883	0,941	1,032	1,985	3,773	> 10
10	1,018	0,903	0,914	0,891	0,947	1,039	1,980	3,765	> 10
11	1,070	0,928	0,932	0,911	0,965	1,062	1,995	3,777	> 10
12	1,120	0,936	0,947	0,918	0,977	1,066	2,009	3,782	> 10

N=100, Limite de Eficiência Semiparamétrico = 0,831, 5000 replicações

que o modelo 1 é particularmente desfavorável, este fato é um forte indício a favor da combinação. O ganho em utilizar a melhor combinação, em relação a Reponderação ou Imputação, é da ordem de 5%, para $N = 100$, e 1% , para $N = 200$, em termos de erro quadrático médio.

Outro ponto interessante é a proximidade entre $\hat{\beta}_{FI}$ e $\hat{\beta}_{RP}$, quando utilizados o mesmos números de termos K e L . Curiosamente, enquanto, na maioria dos casos, o estimador ótimo é $\hat{\beta}_{RP}$, para valores muito altos de K e L , $\hat{\beta}_{FI}$ é mais estável.

A tabela 5 mostra os resultados da simulação com o modelo 2, mais suave, com $N = 200$. Pela similaridade observada, para este modelo e os próximos, são reportados apenas os valores de $\hat{\beta}_{RP}$. Verifica-se que, quando o modelo é muito suave, a diferença entre os todos os estimadores desaparece rapidamente. Nesta especificação, embora o melhor estimador empregue estritamente uma combinação das técnicas, percebe-se que há tanto estimadores de Reponderação quanto de Imputação com precisão distante em menos de 3% do Limite de Eficiência.

A especificação menos suave deteriora a performance de todos os estimadores, conforme mostrado na tabela 6. Percebe-se que os melhores estimadores empregam um grande número de termos na estimação de $p(\cdot)$, e que a introdução de pelo menos um termo além da constante, na estimação de $m_t(\cdot)$, geralmente reduz o erro quadrático. Também é notada uma mudança no desempenho relativo entre Imputação pura e Reponderação pura, favorável à última.

Uma possível explicação para esta observação é que $p(\cdot)$ é uma trans-

Tabela 5.3: Modelo 1, Regressão Ponderada, N=200

		Erros Quadráticos Médios							
K \ L	1	2	3	4	5	6	7	8	9
1	2,557	0,512	0,458	0,439	0,439	0,436	0,445	0,450	0,481
2	0,491	0,542	0,457	0,438	0,438	0,437	0,444	0,450	0,476
3	0,449	0,449	0,459	0,439	0,438	0,436	0,444	0,451	0,476
4	0,440	0,437	0,438	0,440	0,438	0,437	0,444	0,450	0,475
5	0,438	0,433	0,435	0,434	0,438	0,437	0,444	0,450	0,475
6	0,439	0,435	0,436	0,434	0,436	0,437	0,443	0,450	0,476
7	0,443	0,436	0,437	0,435	0,437	0,436	0,444	0,450	0,478
8	0,445	0,435	0,437	0,434	0,436	0,436	0,443	0,451	0,478
9	0,449	0,437	0,439	0,436	0,438	0,437	0,445	0,452	0,478
10	0,450	0,437	0,438	0,435	0,437	0,436	0,445	0,451	0,477
11	0,458	0,439	0,442	0,437	0,440	0,439	0,447	0,453	0,480
12	0,467	0,442	0,444	0,438	0,442	0,440	0,448	0,455	0,483

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

formação de $m_1(.) - m_0(.)$ pela função logística, $L(.)$. Isto pode fazer com que a perda de suavidade tenha sido desigual, pois a composição com $L(.)$ atenua variações.

Os modelos utilizados para avaliar o efeito da heterocedasticidade, embora aparentemente representando casos extremos, não forneceram evidência favorável à sugestão de que esta poderia afetar os benefícios da Regressão Ponderada. As tabelas 7 e 8 reportam os resultados, respectivamente, para os casos supostamente favorável e desfavorável. Os estimadores ótimos, em ambos os casos, utilizam o mesmo número de termos, contrariando a expectativa de que, no segundo, seria recomendada uma reponderação mais precisa.

A introdução de múltiplas variáveis auxiliares reforça o ganho da combinação de estratégias, quando utilizadas pelo estimador de Regressão Ponderada, como mostrado na tabela 9. O melhor estimador Duplamente Robusto, neste caso, tem erro quadrático médio 10% e 8% menor que os melhores estimadores de, respectivamente, Imputação e Reponderação. Diferentemente das especificações anteriores, porém, a forma de combinar os métodos influencia o resultado. Na tabela 10, que considera o estimador de Reponderação Aumentado, ocorrem perdas consideráveis em algumas combinações de K e L , e o melhor estimador utiliza apenas Reponderação. No entanto, o estimador $\hat{\beta}_{FI}$ que utiliza a combinação ótima para $\hat{\beta}_{RP}$ ainda se encontra entre os melhores.

Por fim, conforme esperado, a reponderação pelo inverso do *propensity score* verdadeiro resultou em desvios muito maiores que o Limite de Eficiência Semiparamétrico, como podemos observar na tabela 11. Por outro lado, o uso de $p(.)$ combinado com procedimentos de imputação permitiu atingir desem-

Tabela 5.4: Modelo 1, Reponderação Aumentada, N=200

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	2,557	0,512	0,458	0,439	0,439	0,436	0,445	0,450	0,481
2	0,491	0,548	0,457	0,439	0,439	0,436	0,445	0,450	0,481
3	0,449	0,449	0,461	0,439	0,439	0,436	0,445	0,450	0,481
4	0,440	0,438	0,438	0,441	0,439	0,437	0,445	0,450	0,481
5	0,438	0,433	0,436	0,433	0,439	0,437	0,445	0,451	0,481
6	0,439	0,435	0,437	0,434	0,437	0,437	0,445	0,450	0,481
7	0,443	0,436	0,438	0,435	0,438	0,436	0,445	0,451	0,481
8	0,445	0,436	0,438	0,434	0,437	0,436	0,444	0,451	0,482
9	0,449	0,438	0,440	0,436	0,439	0,437	0,446	0,452	0,481
10	0,450	0,437	0,440	0,434	0,438	0,436	0,445	0,451	0,480
11	0,458	0,440	0,443	0,437	0,441	0,438	0,448	0,453	0,482
12	0,467	0,443	0,447	0,439	0,445	0,440	0,450	0,455	0,484

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

penho semelhante ao das implementações ótimas de $\hat{\beta}_{FI}$ e $\hat{\beta}_{RP}$, exceto nos modelos 3 e 6. Sinteticamente isto é um indício de que o uso de informação sobre o *propensity score*, embora piore o desempenho assintótico em estimadores de Reponderação simples, permite atingir eficiência assintótica quando se combina com a estimação das funções de regressão. Em modelos menos afetados pela ‘maldição da dimensionalidade’, o recurso a esta informação foi particularmente efetivo em facilitar a obtenção de um estimador ótimo, uma vez que eliminou a necessidade de se escolher o número de termos a utilizar em $\hat{p}(\cdot)$, ao mesmo tempo em que possibilitou alcançar um erro quadrático médio praticamente idêntico ao de $\hat{\beta}_{FI}$ e $\hat{\beta}_{RP}$.

Tabela 5.5: Modelo 2, Regressão Ponderada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	2,417	0,442	0,428	0,428	0,430	0,432	0,437	0,443	0,466
2	0,430	0,446	0,428	0,427	0,430	0,432	0,437	0,443	0,467
3	0,429	0,427	0,428	0,428	0,430	0,432	0,437	0,443	0,467
4	0,430	0,427	0,427	0,428	0,429	0,432	0,437	0,443	0,468
5	0,433	0,429	0,429	0,429	0,429	0,433	0,437	0,443	0,468
6	0,438	0,430	0,431	0,430	0,431	0,433	0,437	0,444	0,467
7	0,448	0,431	0,432	0,432	0,433	0,434	0,438	0,444	0,467
8	0,446	0,433	0,433	0,433	0,433	0,435	0,438	0,444	0,467
9	0,452	0,434	0,434	0,433	0,434	0,436	0,439	0,445	0,468
10	0,456	0,436	0,437	0,436	0,437	0,438	0,442	0,447	0,469
11	0,475	0,440	0,440	0,439	0,440	0,441	0,445	0,449	0,472
12	0,474	0,442	0,443	0,442	0,443	0,443	0,447	0,450	0,474

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

Tabela 5.6: Modelo 3, Regressão Ponderada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	2,683	1,346	0,892	0,674	0,598	0,533	0,532	0,503	0,548
2	1,357	1,435	0,908	0,678	0,595	0,530	0,524	0,495	0,524
3	0,879	0,900	0,950	0,703	0,600	0,531	0,524	0,495	0,524
4	0,679	0,685	0,702	0,734	0,610	0,536	0,520	0,492	0,520
5	0,578	0,576	0,586	0,595	0,627	0,545	0,525	0,496	0,521
6	0,523	0,520	0,523	0,526	0,544	0,556	0,532	0,497	0,525
7	0,491	0,486	0,491	0,489	0,506	0,505	0,538	0,501	0,527
8	0,479	0,470	0,473	0,471	0,482	0,479	0,504	0,504	0,533
9	0,458	0,453	0,462	0,457	0,469	0,461	0,486	0,479	0,537
10	0,449	0,448	0,453	0,449	0,457	0,452	0,469	0,468	0,512
11	0,446	0,442	0,446	0,440	0,448	0,441	0,459	0,456	0,499
12	0,451	0,445	0,450	0,442	0,452	0,444	0,464	0,459	0,504

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

Tabela 5.7: Modelo 4, Regressão Ponderada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	2,620	0,531	0,457	0,445	0,440	0,438	0,450	0,448	0,477
2	0,504	0,562	0,458	0,443	0,439	0,438	0,448	0,448	0,471
3	0,450	0,451	0,461	0,445	0,439	0,438	0,448	0,448	0,470
4	0,444	0,441	0,442	0,445	0,439	0,438	0,447	0,448	0,470
5	0,440	0,436	0,437	0,437	0,440	0,438	0,447	0,448	0,471
6	0,443	0,436	0,437	0,436	0,438	0,438	0,447	0,448	0,473
7	0,446	0,438	0,440	0,438	0,440	0,440	0,447	0,448	0,473
8	0,448	0,438	0,440	0,437	0,440	0,439	0,446	0,449	0,476
9	0,451	0,440	0,442	0,439	0,441	0,440	0,447	0,450	0,477
10	0,457	0,443	0,443	0,440	0,443	0,442	0,449	0,451	0,480
11	0,462	0,445	0,445	0,442	0,445	0,443	0,451	0,453	0,482
12	0,466	0,446	0,446	0,443	0,445	0,444	0,451	0,453	0,483

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

Tabela 5.8: Modelo 5, Regressão Ponderada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	2,564	0,521	0,453	0,432	0,431	0,431	0,447	0,453	0,483
2	0,494	0,550	0,454	0,432	0,431	0,432	0,446	0,452	0,472
3	0,443	0,446	0,457	0,433	0,431	0,432	0,446	0,455	0,471
4	0,432	0,432	0,432	0,435	0,431	0,432	0,446	0,455	0,473
5	0,432	0,430	0,430	0,430	0,432	0,433	0,447	0,455	0,473
6	0,434	0,430	0,431	0,430	0,431	0,434	0,447	0,455	0,474
7	0,441	0,433	0,434	0,432	0,434	0,435	0,447	0,456	0,474
8	0,443	0,435	0,436	0,434	0,436	0,437	0,449	0,458	0,474
9	0,450	0,437	0,439	0,436	0,438	0,439	0,452	0,460	0,475
10	0,457	0,442	0,443	0,440	0,442	0,443	0,455	0,463	0,477
11	0,459	0,442	0,445	0,441	0,444	0,444	0,457	0,464	0,477
12	0,465	0,444	0,447	0,443	0,445	0,445	0,459	0,465	0,479

N=200, Limite de Eficiência Semiparamétrico = 0,4155, 5000 replicações

Tabela 5.9: Modelo 6, Regressão Ponderada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	6,313	5,336	4,547	4,250	4,201	4,463	4,519	3,826	3,906
2	5,319	5,280	4,520	4,184	4,093	4,349	4,393	3,791	3,892
3	4,443	4,505	4,478	4,160	4,089	4,183	4,234	3,738	3,803
4	3,995	4,068	4,026	4,178	4,094	4,164	4,189	3,677	3,760
5	3,878	3,967	3,925	4,078	4,100	4,172	4,195	3,671	3,763
6	3,805	3,892	3,928	4,093	4,090	4,189	4,208	3,695	3,783
7	3,825	3,899	3,940	4,089	4,075	4,154	4,197	3,668	3,768
8	3,746	3,725	3,630	3,785	3,762	3,836	3,887	3,673	3,779
9	3,755	3,727	3,622	3,750	3,728	3,796	3,868	3,636	3,758
10	4,124	3,910	3,696	3,669	3,640	3,720	3,779	3,475	3,600
11	4,006	3,844	3,657	3,663	3,636	3,726	3,791	3,513	3,625
12	3,851	3,847	3,738	3,744	3,712	3,796	3,862	3,547	3,669

N=100, 5000 replicações

Tabela 5.10: Modelo 6, Reponderação Aumentada

Erros Quadráticos Médios Simulados									
K \ L	1	2	3	4	5	6	7	8	9
1	6,313	5,336	4,547	4,250	4,201	4,463	4,519	3,826	3,906
2	5,319	5,286	4,574	4,295	4,200	4,462	4,519	3,827	3,907
3	4,443	4,547	4,536	4,281	4,212	4,312	4,363	3,832	3,907
4	3,995	4,113	4,096	4,280	4,203	4,275	4,304	3,760	3,870
5	3,878	4,039	4,026	4,218	4,238	4,315	4,339	3,766	3,894
6	3,805	3,979	4,044	4,268	4,260	4,362	4,377	3,806	3,933
7	3,825	3,993	4,052	4,267	4,262	4,330	4,364	3,778	3,899
8	3,746	3,824	3,765	4,037	4,024	4,067	4,127	3,847	3,986
9	3,755	3,848	3,799	4,045	4,032	4,066	4,149	3,851	3,999
10	4,124	4,319	4,151	4,460	4,430	4,361	4,473	3,875	4,049
11	4,006	4,183	4,047	4,407	4,399	4,386	4,474	3,925	4,072
12	3,851	3,999	3,958	4,203	4,214	4,321	4,313	3,823	3,961

N=100, 5000 replicações

Tabela 5.11: Estimativas usando *propensity score* verdadeiro

Erros Quadráticos Médios Simulados								
Modelo	N \ L	1	2	3	4	6	8	9
1, $\tilde{\beta}_{RP}$	100	1,348	0,906	0,870	0,856	0,992	3,899	> 10
1, $\tilde{\beta}_{RP}$	200	0,661	0,459	0,442	0,430	0,434	0,448	0,502
2, $\tilde{\beta}_{RP}$	200	0,741	0,436	0,428	0,428	0,436	0,463	0,506
3, $\tilde{\beta}_{RP}$	200	0,510	0,494	0,501	0,491	0,484	0,484	0,789
4, $\tilde{\beta}_{RP}$	200	0,709	0,471	0,447	0,437	0,438	0,447	0,488
5, $\tilde{\beta}_{RP}$	200	0,656	0,452	0,436	0,429	0,432	0,452	0,484
6, $\tilde{\beta}_{RP}$	100	4,580	4,633	4,573	4,299	4,427	3,853	3,873
6, $\tilde{\beta}_{FI}$	100	4,580	4,687	4,649	4,408	4,587	3,946	3,924