

5

Metodologia para avaliação e mensuração do risco moral e seleção adversa

Nesta dissertação procurar-se-á avaliar o risco moral e a seleção adversa por intermédio de análise estatística. Espera-se poder evidenciar se tanto o risco moral quanto a seleção adversa estão presentes ou não no mercado de planos de saúde do Brasil.

Com relação à seleção adversa se propõe um estudo onde é ajustado um modelo logístico para analisar a relação de algumas variáveis com o fato de um indivíduo ter plano de saúde, tais como sexo, renda, idade, grau de instrução, etc.

A busca pela avaliação do risco moral está estruturada em três etapas: na primeira parte são feitas análises exploratórias a partir de informações do Suplemento Saúde da PNAD 98. Esta primeira etapa terá por objetivo verificar estatísticas que venham a informar se há indícios da presença do risco moral no cenário brasileiro de planos de saúde.

Na segunda parte desta sequência se propõe a construção do indicador de risco moral, o IRM. A diferença marcante do IRM para as informações preliminares levantadas na primeira etapa com a análise exploratória é que para construção do IRM serão levados em consideração os efeitos da incorporação do plano amostral da PNAD para obtenção de intervalos de confiança que possibilitem a avaliação dos números médios de consultas médicas realizados por segurados com (ou sem) planos de saúde. É a diferença estatística significativa entre o número de consultas para os dois grupos que indicará a presença do risco moral. Como na estimativa do IRM não são considerados efeitos de variáveis que podem influir nas probabilidades de uma pessoa realizar consulta médica, como renda, sexo, idade estima-se então em uma terceira etapa um modelo econométrico hurdle binomial negativo a fim de que se alcance um resultado mais contundente sobre o número médio de consultas de pessoas com e sem plano de saúde.

5.1.

Análise exploratória dos dados da PNAD/IBGE

O Suplemento Saúde da PNAD 98 constitui-se num instrumento importante para análise das características populacionais relacionadas ao consumo de serviços de saúde no Brasil. Nesta parte da dissertação, a partir das informações contidas neste suplemento, serão apresentados os principais resultados da pesquisa que, além de fornecerem insumos para conhecimento do retrato brasileiro na área de saúde, também permitem um levantamento inicial para observação de evidências empíricas sobre o risco moral no âmbito do nosso estudo, ou seja, no mercado brasileiro de planos de saúde.

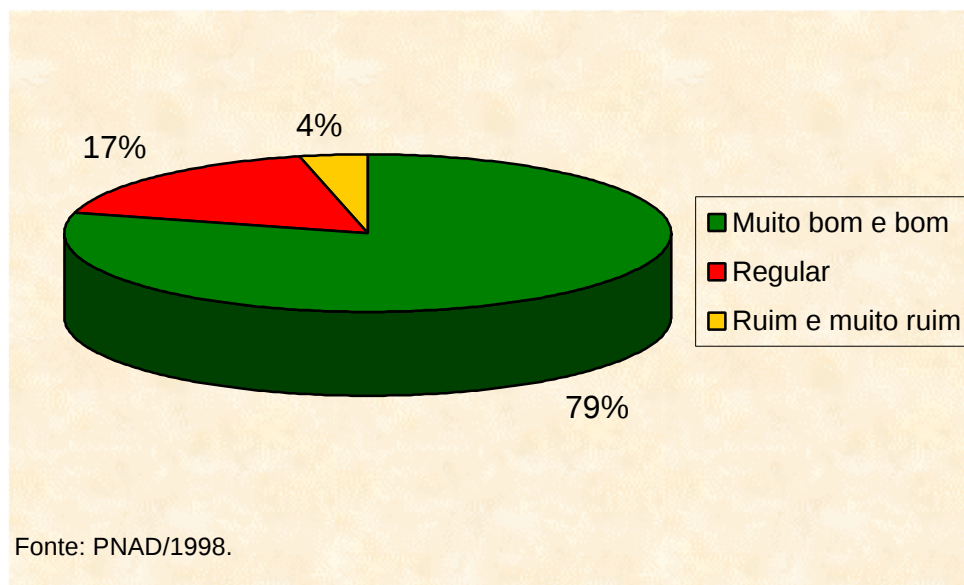
Esta parte de exploração de dados compreenderá basicamente duas fases: uma relacionada às características de saúde dos brasileiros; outra focando dados que possam caracterizar a presença de informação assimétrica no mercado brasileiro e suas consequências principais, o risco moral e seleção adversa. Esta análise exploratória inicial será a base do restante da dissertação, já que são as evidências aqui encontradas que pautarão o desenvolvimento do restante do trabalho. Na conclusão da análise exploratória são incluídos os argumentos e as justificativas para os desdobramentos e desenvolvimento das etapas posteriores do estudo.

5.1.1.

Algumas características de saúde dos brasileiros

Em 1998, data de referência da PNAD 98, a população brasileira compreendia cerca de 158 milhões de habitantes. Deste total, 79,1% da população considerava o próprio estado de saúde como sendo bom ou muito bom; 17,2% se auto-avaliaram como tendo um estado de saúde regular e apenas 3,6% dos brasileiros informaram considerar o próprio estado de saúde como sendo ruim ou muito ruim. A figura 1 ilustra o conteúdo destes comentários:

Figura 1 – Distribuição percentual da população residente, segundo auto-avaliação do estado de saúde.



Segundo as notas metodológicas da PNAD 98, para efeito de análise, denominou-se índice de satisfação ao percentual de indivíduos, em cada categoria, que declararam o estado de saúde como sendo bom ou muito bom. A partir deste conceito algumas conclusões foram obtidas. Por exemplo, o número de homens que se consideraram como tendo um estado de saúde bom ou muito bom é maior do que o número de mulheres com a mesma classificação do estado de saúde. O índice de satisfação dos homens foi estimado em 81,8%. Já para as mulheres o índice ficou em 76,4%.

Com foco na faixa etária, se verifica que para pessoas com 14 anos ou menos de idade o índice é muito semelhante quando são comparados os dois sexos. Neste caso o índice tanto para os homens quanto para as mulheres foi de 92,0%.

Porém, para as demais faixas etárias, o índice de satisfação das mulheres é sempre menor que o dos homens. Os valores mais baixos do índice de satisfação foram para as pessoas que, proporcionalmente, utilizam mais os serviços de saúde, ou seja, pessoas com idade superior a 64 anos. Para este grupo de indivíduos o índice de satisfação é muito baixo. Para as mulheres é de 34,2% e para os homens é de 39,4%.

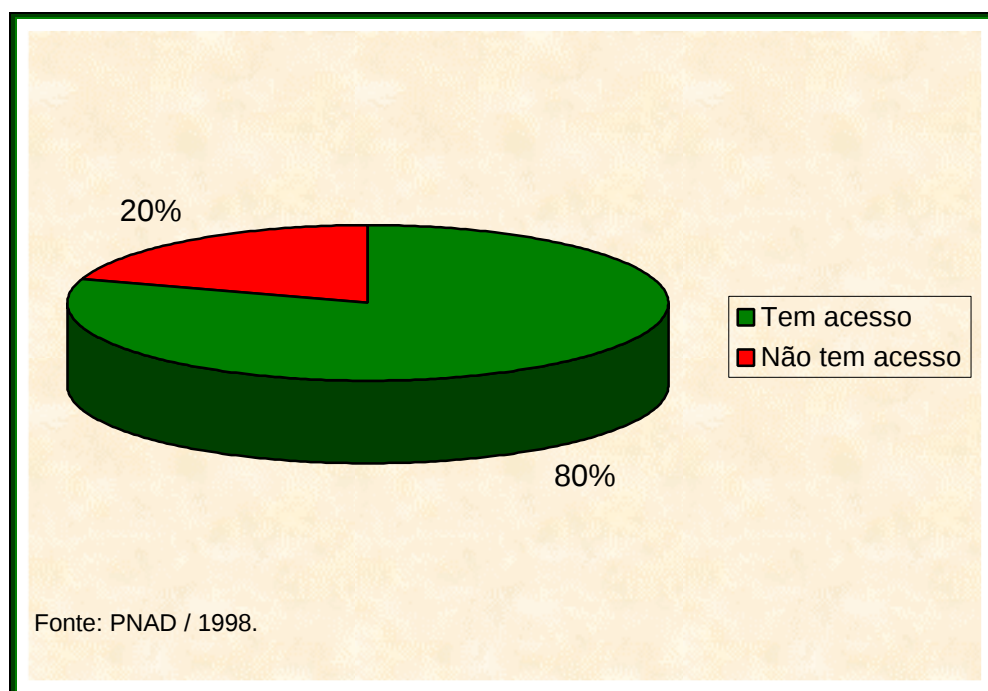
Relativamente às doenças crônicas, as informações oriundas da PNAD dão conta de que 31,6% da população brasileira sofre de pelo menos um tipo de doença crônica. E se observa que as mulheres são as que mais sofrem, já que

35,3% têm pelo menos uma doença crônica enquanto que para os homens este valor cai para 27,7%.

No Brasil ainda existem pessoas idosas que não têm acesso a serviços de saúde e a rede de atendimento público precisa ser ampliada para suportar essa demanda. Deve-se ainda considerar que quando se diz que o acesso à saúde deve ser oferecido pelo governo às pessoas idosas não se está excluindo as pessoas que atualmente estão cobertas por planos de saúde. Se o problema de seleção adversa existir, esta é uma medida que deve ser tomada para eliminá-lo.

Como pode ser observado na figura 2, mesmo sem considerar as pessoas cobertas por algum seguro-saúde, ainda há carência na rede de postos de atendimento, já que se estima que 20% da população com mais de 65 anos de idade não têm acesso a qualquer serviço de saúde.

Figura 2 – Acesso aos serviços de saúde para pessoas com mais de 65 anos.



As principais conclusões da PNAD 98 relacionadas à saúde dos brasileiros, derivadas de informações baseadas na auto-avaliação do estado de saúde e restrição de atividades habituais por motivo de saúde ou doença crônica foram que as necessidades em saúde têm um padrão de distribuição, segundo a idade, ou seja, as pessoas no início e particularmente no final da vida, apresentam mais problemas de saúde. Os homens referem mais problemas de saúde do que as

mulheres apenas nas idades mais jovens - início da adolescência. A partir dos 14 anos de idade são as mulheres que passam a referir problemas de saúde com maior frequência. Também foi possível observar que estes padrões referentes à idade e sexo são semelhantes aos observados em outros países.

Também como em outros países, observa-se no Brasil que o número de pessoas que referem problemas de saúde diminui à medida em que a renda familiar aumenta, definindo um padrão de marcadas desigualdades sociais em saúde e mostrando, mais uma vez, que a desigualdade social no Brasil é um problema muito sério e que precisa ser combatido, a fim de que a população, em diversos campos, possa ter um nível melhor de qualidade de vida.

5.1.2.

Estatísticas descritivas que permitem uma primeira avaliação do risco moral e da seleção adversa

Nesta parte do trabalho é dada seqüência a análise exploratória. Serão comentados os resultados das estatísticas descritivas da PNAD 98 que têm relação com o estudo do risco moral e da seleção adversa.

No capítulo 4 desta dissertação, quando foram explicitados os conceitos de risco moral e seleção adversa, argumentou-se que na área de saúde o risco moral se caracteriza pela maior utilização dos serviços de saúde por parte de um segurado em função deste segurado estar coberto por plano de saúde. Por exemplo, se verifica o risco moral se uma pessoa realiza mais visitas médicas do que faria se não tivesse plano. Logo, se espera com a análise de dados verificar se realmente as pessoas com planos de saúde realizam mais consultas do que as pessoas não cobertas. Evidentemente, esta análise exploratória é limitada pela não possibilidade de controle de outras variáveis que possam influenciar as conclusões sobre risco moral. E é exatamente por isso que é considerada como sendo apenas uma etapa para levantamento de informações iniciais que, somadas às provenientes das outras etapas do desenvolvimento metodológico, permitirão, espera-se, uma investigação mais completa sobre os fenômenos estudados.

A seleção adversa, por sua vez, fica bem definida quando as seguradoras estão diante de uma situação desfavorável, representada, por exemplo, pelo fato de muitas pessoas de alto risco estarem aderindo a um determinado tipo de seguro

que não considera, na cobrança do prêmio, este incremento do risco, fazendo com que as seguradoras possam vir, em casos mais graves, a se encontrarem em uma situação de falência.

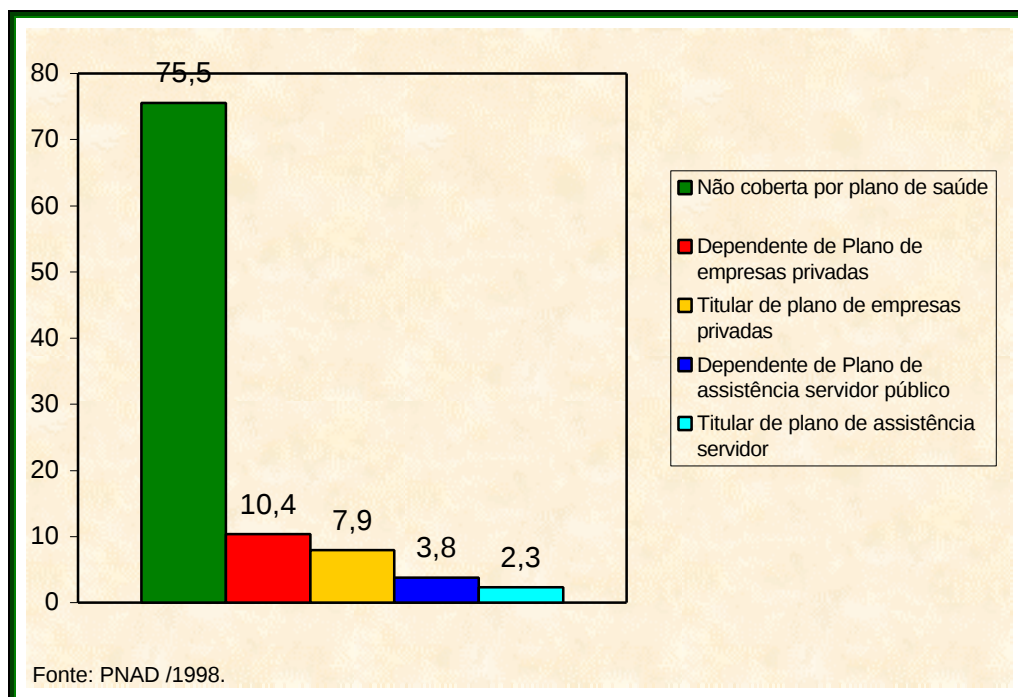
No âmbito deste estudo a seleção adversa pode ficar bem caracterizada pela verificação da situação em que os planos de saúde cubram muitas pessoas de risco elevado, pelo menos em maior número que pessoas de baixo risco. Na verdade, ainda para se assegurar a ocorrência de seleção adversa no mercado brasileiro seria preciso assumir que não existem mecanismos de defesa por parte das seguradoras, como a diferenciação dos preços das mensalidades por faixas etárias. Porém, a diferenciação existe, fazendo com que a confirmação da ocorrência da seleção adversa possa ser restringida a avaliação das participações e proporções de indivíduos com alto risco ou baixo risco nos planos de saúde. A análise descritiva dos dados da PNAD 98 pode ser importante para subsidiar a análise proposta.

5.1.2.1.

Cobertura por plano de saúde

Segundo a PNAD 98, 38,7 milhões de pessoas no Brasil tem cobertura de planos de saúde. Este percentual, levado para o total da população, corresponde a 24,5% dos brasileiros. Deste total, 29 milhões (75% dos que têm plano) estão vinculados a planos de saúde privados e 9,7 milhões (25%) estão vinculados a planos de instituto ou instituição patronal de assistência ao servidor público civil ou militar. Com relação à estratificação urbana e rural percebe-se que a cobertura de planos de saúde é absolutamente maior nas áreas urbanas, com percentual de cobertura chegando a 29,2% enquanto que nas áreas rurais do país o percentual de cobertura é de apenas 5,8%. A Figura 3 ilustra a distribuição percentual da população residente, segundo a cobertura de planos de saúde.

Figura 3 – Distribuição percentual residente, segundo a cobertura de planos de saúde.



Quando se analisa a distribuição por sexo verifica-se que 25,7% das mulheres estão cobertas por planos enquanto que no caso dos homens a cobertura chega a 23,0%. Como era de se esperar, o percentual de indivíduos que possuem planos de saúde por faixa etária varia. Para pessoas com idade de até 18 anos a cobertura por plano atinge cerca de 20,7%. Para o grupo de pessoas com idades entre 40 e 64 anos o percentual de cobertura é de 29,5%. Para as pessoas mais idosas, com idades acima dos 64 anos a cobertura cai para 26,1% para os homens e 28,2% para as mulheres.

É possível também notar uma associação positiva entre cobertura de plano de saúde e renda familiar: a cobertura é de 2,6% na classe de renda familiar inferior a um salário mínimo. Quando se passa para pessoas com faixa de renda familiar compreendida entre 1 e 2 salários mínimos a cobertura cresce para 4,8%, e continua a crescer com maior intensidade nas demais classes de renda: 9,4% (2 a 3 salários mínimos), 18,0% (3 a 5 salários mínimos), 34,7% (5 a 10 salários mínimos) e 76% (20 salários mínimos e mais), como mostra a figura 4, a seguir:

Figura 4 – População residente, por cobertura de plano de saúde e classes de rendimento familiar.

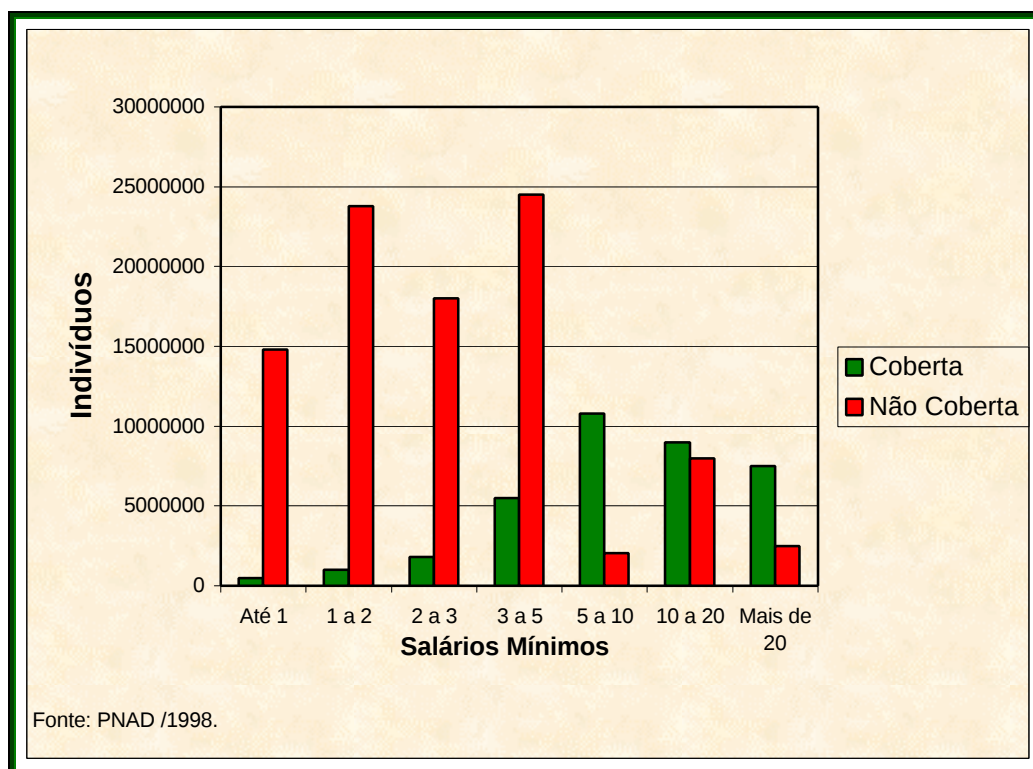


Tabela 1 – Auto-avaliação de saúde por cobertura de plano

Auto avaliação de saúde	Com plano
Muito bom e bom	83,6%
Regular	14,1%
Ruim e muito ruim	2,1%
Sem avaliação	0,2%

Fonte: PNAD/1998.

Dentre as pessoas que têm plano de saúde, 83,6% avaliaram o próprio estado de saúde física e mental como sendo bom ou muito bom. 14,1% das pessoas com plano avaliaram o próprio estado de saúde como regular e apenas 2,1% como sendo ruim ou muito ruim.

Esta última informação, resumida na tabela 1 é muito importante para a questão da seleção adversa. Se somente esta análise descritiva é considerada, nota-se indícios de que, no mercado de planos de saúde brasileiro, a seleção

desfavorável não ocorre por parte dos segurados com relação às administradoras de planos de saúde.

Esta suposição emerge do fato de que, segundo as evidências da PNAD 98, a cobertura por plano de saúde é maior entre as pessoas que avaliam o próprio estado de saúde como “bom ou muito bom” e diminui na medida em que a auto-avaliação do estado de saúde piora. Este quadro em valores percentuais sinaliza que, ao contrário do que se espera para caracterização da seleção adversa, o mercado de planos no Brasil não apresenta indícios de uma situação desfavorável para as seguradoras e administradoras de planos de saúde, já que a maior parte das pessoas cobertas por plano estão em bom estado de saúde. Porém, para uma afirmação mais contundente, ainda é preciso conhecer esses dados em números absolutos.

E, com a análise da tabela 2, abaixo, fica realmente claro que, se a auto avaliação do estado de saúde pode ser assumida como proxy do estado de saúde, então há fortes indícios de que não há uma situação desfavorável para as seguradoras no mercado brasileiro de planos de saúde, pois a maior parte das pessoas com planos de saúde efetivamente pertence ao grupo das pessoas que classificaram o próprio estado de saúde como sendo “bom ou muito bom”. São mais de 32 milhões de indivíduos em bom estado de saúde contra apenas cerca de 836 mil pessoas em estado ruim.

Tabela 2 – Auto-avaliação do estado de saúde de pessoas cobertas por planos de saúde em números absolutos.

Estado de saúde	Em milhões
Muito bom e Bom	32 363 730
Regular	5 471 160
Ruim e muito ruim	836 221
Sem declaração	9 295

Fonte: PNAD/1998.

5.1.2.2.

Utilização dos serviços de saúde

Dando prosseguimento à análise descritiva dos dados da PNAD 98, agora com relação ao acesso e utilização dos serviços de saúde, encontra-se, ao contrário do que se verificou para seleção adversa, indicativos de que o risco moral está presente no mercado brasileiro.

Neste âmbito de análise observa-se que dentre as pessoas que procuram regularmente os Postos ou Centros de Saúde não existe uma diferenciação significativa segundo sexo, já que apresentam valores similares próximos a 42%. Em função da idade observou-se que a população jovem, com idade até 18 anos, é a que mais procura este tipo de serviço. À medida que aumenta a renda familiar mensal, menor é a procura pelos Postos e Centros de Saúde. Cerca de 55,0% da população nas duas primeiras faixas de renda média familiar (até 2 salários mínimos) declararam procurar, em caso de necessidade, serviços com estas características.

As pessoas que disseram procurar atendimento em ambulatórios de Hospitais não apresentaram nenhuma diferença segundo o sexo e idade, constatando-se porcentagens próximas aos 21,0%, nos homens assim como nas mulheres, e em todas as faixas etárias. Esta modalidade de serviço médico é procurada, majoritariamente, pela população com menor nível de renda (até dois salários mínimos).

A população idosa, em especial as mulheres com maior nível de renda, é a que procura com maior intensidade os consultórios particulares quando necessita de algum tipo de atendimento à saúde. As farmácias são mais procuradas pelos homens adultos jovens, de 19 a 39 anos, e pela população que declarou até 1 salário mínimo de renda média mensal.

Tabela 3 – Procura por consultas médicas nos últimos 12 meses por grupos de indivíduos

Mulheres	62,30%
Crianças com menos de 4 anos	68,40%
Maiores de 65 anos	73,20%

Fonte: PNAD/1998.

Com relação a procura por serviços médicos, segundo a PNAD 98 aproximadamente 86,5 milhões de pessoas (54,7% da população brasileira) declararam ter consultado médico nos últimos 12 meses. De acordo com a tabela 3 os grupos populacionais que mais consultaram médico no ano anterior foram: mulheres (62,3%); crianças menores de 4 anos (68,4%); pessoas maiores de 65 anos (73,2%); e os residentes em áreas urbanas (57,2%).

Um dado importantíssimo para esta dissertação, entretanto, surge da observação da alta correlação positiva entre acesso aos serviços médicos e o poder aquisitivo da população. Aproximadamente 49,7% das pessoas de menor renda familiar declararam ter realizado consulta médica nos últimos 12 meses. Esse valor sobe para 67,2% no caso daquelas pessoas com mais de 20 salários mínimos de renda familiar.

Quando esta informação é cruzada com a informação de que as pessoas com maior poder aquisitivo são as que têm mais cobertura de planos de saúde e ainda quando se admite a hipótese de que um maior poder aquisitivo não determina uma maior probabilidade de enfermidade, então se pode pensar, num primeiro momento, que pessoas com planos de saúde estão se consultando mais do que pessoas sem planos de saúde. Isto se constitui num primeiro indício do risco moral no mercado brasileiro.

Este indício é reforçado quando da análise de um outro resultado da PNAD 98, que diz que mais de 7 milhões de consultas foram realizadas por pessoas cobertas por planos num período de duas semanas anteriores a data da entrevista, representando 19% do total de indivíduos com plano de saúde, enquanto que para o grupo de indivíduos não cobertos o percentual de consultas cai para 10%, cerca de duas vezes menos.

5.1.3.

Algumas importantes observações da análise exploratória de dados

Com base nos comentários descritos anteriormente, pode-se dizer que o artifício da diferenciação de preços em função da faixa etária, empregado hoje no Brasil, parece, a princípio, estar funcionando de forma satisfatória no sentido de evitar uma situação adversa para as seguradoras. Entretanto, uma análise mais específica é necessária. Para o risco moral, por outro lado, há fortes indícios de que o mesmo esteja ocorrendo de forma significativa no Brasil.

Em decorrência desta conclusão esta dissertação focará a seleção adversa e o risco moral de forma mais detalhada a partir desta etapa. Com relação à seleção adversa, será realizado um estudo econométrico com o intuito de aprofundar as conclusões sobre o tema. Para o risco moral será construído o indicador de risco moral, denominado IRM. Complementando a análise, também se propõe, para o caso do risco moral, inferências a partir da estimação de um modelo econométrico denominado de hurdle binomial negativo. Este estudo econométrico será detalhado mais adiante.

5.2.

Avaliação econométrica da seleção adversa

Na seção 5.1 a análise exploratória de dados possibilitou a identificação de estatísticas e valores que indicaram preliminarmente a não ocorrência da seleção adversa no mercado de planos de saúde. Esta conclusão se deu, principalmente, pela observação de que a maior parte das pessoas que tinham planos de saúde eram pessoas em bom estado de saúde física e mental. Se a suposição de que a maioria das pessoas pode conhecer o próprio estado saúde não for absurda, então pode-se inferir que as administradoras de planos de saúde têm suas carteiras de clientes compostas em grande parte, por pessoas saudáveis. Esta verificação é respaldada pela observação de que um pequeno número de pessoas idosas possui cobertura de plano de saúde.

Entretanto, o estudo descritivo de dados não levou em consideração possíveis influências de outras variáveis nos resultados. Por isso, nesta parte do trabalho, busca-se a avaliação da seleção adversa de uma forma mais global, por intermédio da identificação de fatores que determinam a aquisição de planos de

saúde. Por exemplo, se pessoas de maior risco adquirem mais planos com cobertura completa do que pessoas de baixo risco, as seguradoras precisam continuar com as políticas de diferenciação de preços (até certos limites), a fim de que não entrem em situação de falência. Evidentemente, o quão diferenciados devem ser os preços das mensalidades ou dos prêmios em função dos riscos depende, dentre outros fatores, da magnitude dos retornos esperados que as seguradoras estariam dispostas a receber, mas esta análise foge aos propósitos desta dissertação.

Mais uma vez a análise se baseará na PNAD 98. Espera-se verificar se as administradoras estão sujeitas a uma situação desfavorável em função de variáveis que determinam o fato de um indivíduo ter ou não ter plano de saúde. Porém, antes da análise econométrica, considera-se importante descrição de aspectos relacionados às características das pesquisas e dados analisados em estudos econométricos na área de saúde. Estes aspectos devem ser observados a fim de que as limitações das técnicas utilizadas possam ser conhecidas e avaliadas. Complementando os tópicos seguintes, também serão feitos comentários sobre o plano amostral da PNAD 98.

5.2.1. Incorporação do plano amostral da PNAD e seus efeitos

As pesquisas amostrais constituem importante fonte de informação para propósitos descritivos e analíticos e oferecem uma série de vantagens quando comparadas aos censos populacionais. Os conceitos fundamentais da amostragem estatística foram estabelecidos por Neyman, em 1934, que introduziu o conceito de estratificação e seleção com probabilidades desiguais. A partir de então outros métodos de amostragem foram desenvolvidos, sendo que a técnica estatística mais utilizada pelos amostristas até os dias de hoje, é a amostragem proporcional ao tamanho, desenvolvida por Horwitz e Thompson em 1952.

É denominado de plano amostral complexo aquele que em seu desenho foram incorporadas alguns níveis de complexidade, tais como estratificação, conglomeração e probabilidades desiguais de seleção. Dependendo de como é

incorporado o plano amostral na estimação de medidas descritivas, é possível obter estatísticas não viciadas com respeito ao desenho amostral.

Geralmente, os pesquisadores assumem que os dados foram coletados a partir de uma amostra aleatória simples, onde as observações são independentes e identicamente distribuídas. Ou então utilizam os dados como se fossem informações coletadas a partir de um censo (sem utilizar os pesos). Os grandes institutos produtores de estatísticas normalmente incorporam os pesos e a complexidade dos dados em suas análises e produzem resultados ajustados com relação a esses fatores. Porém, fora desses institutos, são comuns casos em que pesquisadores não levam em conta o plano amostral e a amostragem, produzindo desta forma, estatísticas viciadas.

Os *softwares* estatísticos, em geral, não possuem ferramentas que incorporem a complexidade de um plano amostral em sua totalidade. Assim as estimativas pontuais dos parâmetros são influenciadas pela ocorrência de pesos distintos, enquanto que as estimativas das variâncias dos estimadores dos parâmetros do modelo são influenciadas pelos efeitos de estratificação e conglomeração.

Todos esses problemas podem ser graves para os estudos estatísticos e econométricos. O efeito de se ignorar o plano amostral pode causar superestimação nas estimativas pontuais de estatísticas descritivas, mais precisamente da média (quando os pesos amostrais não são utilizados). Também se verificam diferenças entre as estimativas do erro padrão dos parâmetros dos modelos ajustados, obtendo assim intervalos de confiança maiores que os “verdadeiros”.

Nesta dissertação, propõe-se a construção do indicador de risco moral – IRM, a partir de dados da PNAD 98, pesquisa que possui um plano amostral complexo. Devido as considerações aqui abordadas, irá se considerar este desenho amostral complexo utilizando-se, para isso, o *software* SUDAAN⁵. Este software oferece a opção do “Método de Linearização de Taylor” na estimação do erro padrão dos estimadores de Máxima Pseudo-Verossimilhança dos parâmetros dos modelos de regressão ajustados. Posteriormente esta metodologia será descrita.

⁵ O SUDAAN foi desenvolvido pelo reseach triangle institute, 1997.

O SUDAAN oferece um tipo de estratificação conhecido como WR⁶, sendo considerado como o mais adequado para aplicação nas análises da PNAD. Este método, assim como outras opções para a definição do desenho amostral consideradas no SUDAAN, estão detalhadas em Vieira (2001).

Com relação aos efeitos do plano amostral, em 1965 Kish desenvolveu um método para comparar ganhos ou perda de precisão sob diferentes planos amostrais no estágio de planejamento da pesquisa. Esta medida foi denominada de Efeito do Plano Amostral (EPA ou DEFF – *Design Effect*). O EPA equivale a razão entre a variância $(V_{verd}(\hat{b}))$ de um estimador verdadeiro, obtido sob a consideração do plano amostral complexo, e a variância $(V_{AAS}(\hat{b}))$ do estimador que seria estimada com base em um plano de amostragem aleatória simples. A expressão do EPA é, então, a seguinte:

$$EPA_{Kish}(\hat{b}) = \frac{V_{verd}(\hat{b})}{V_{AAS}(\hat{b})}$$

Entretanto, os valores de EPA de Kish calculados têm pouca utilidade para estudos analíticos, pois este procedimento se dá no estágio de planejamento da pesquisa. Uma vez que a amostra já esteja selecionada, não há o porque de compará-la com outros planos amostrais.

A partir destas observações surgiu a necessidade de uma nova medida, que teria por finalidade avaliar a tendência de um estimador consistente, calculado sob hipótese de IID, subestimar ou superestimar a variância verdadeira do estimador pontual. Foi aí que em 1989, Skinner, Holt e Smith desenvolveram o EPA ampliado, denominado por *misspecification effect* (meff). Esta medida é capaz de mensurar os efeitos de especificação incorreta tanto do plano amostral quanto do modelo ajustado. O EPA ampliado tem a expressão:

$$EPA(\hat{b}, v_0) = \frac{V_{verd}(\hat{b})}{E_{verd}(v_0)}$$

onde:

$v_0 = \hat{V}_{IID}(\hat{b})$ é um estimador consistente da variância de $\hat{b}$, assumindo-se a hipótese IID para as observações;

⁶ Nascimento Silva recomenda a escolha desta opção para se trabalhar com as bases de dados do IBGE.

$V_{verd}(\hat{b})$ é a variância do estimador sobre o plano utilizado;

$E_{verd}(v_0)$ é a expectância do estimador consistente sob o plano amostral efetivamente utilizado.

A partir dos valores encontrados para o EPA pode-se tirar as seguintes conclusões:

Se o EPA é menor que 1 então a variância sob amostragem aleatória simples é superestimada; Se o EPA é igual a 1 então se conclui que não há evidências de diferença entre as estimativas das variâncias verdadeiras e, por fim, se o EPA é maior do que 1 pode-se concluir que a variância, sob a hipótese de amostragem aleatória simples, está subestimada.

Exemplos do cálculo de EPA de Kish e $EPA(\hat{b}; v_0)$ podem ser encontrados em Pessoa e Nascimento Silva (1998).

A partir da estimativa pontual $\hat{b}$ de um parâmetro b , pode-se construir um intervalo de confiança aproximado $(1-\alpha)$ para o parâmetro com base na distribuição assintótica de: $t_0 = \frac{\hat{b} - b}{\sqrt{v_0}}$, sob hipótese das observações serem IID que tem distribuição Normal padrão, ou seja, $t_0 \underset{IID}{\sim} N(0,1)$. Sob estas condições um Intervalo de Confiança aproximado ao nível de confiança $(1-\alpha)$ é dado por:

$$\left[\hat{b} - Z_{\alpha/2} \sqrt{v_0}; \hat{b} + Z_{\alpha/2} \sqrt{v_0} \right]$$

Onde $Z_{\alpha/2}$ é definido por $\int_{Z_{\alpha/2}}^{\infty} j(t) dt = \frac{\alpha}{2}$ e j é a função de densidade da

distribuição Normal padrão.

Particularmente, nesta dissertação, é este o intervalo de confiança que será usado para inferências sobre diferenças no IRM.

Porém, a distribuição assintótica de t_0 , sob o verdadeiro plano amostral, não é Normal padrão. Mas fazendo a especificação correta do plano amostral temos que a distribuição assintótica verdadeira para t_0 é dada por $t_0 \sim N(0; EPA(\hat{b}, v_0))$.

Analisando-se o valor do $EPA(\hat{b}; v_0)$, pode-se encontrar Intervalos de Confiança baseados na premissa das observações serem IID, diferentes daqueles onde se considera a verdadeira distribuição de t_0 .

Vários valores de EPAs foram calculados e podem ser vistos na tabela abaixo:

Tabela 4 – Níveis de significância e reais

$EPA \quad (\hat{b}, v_0)$	$1 - \alpha = 0,95$	$1 - \alpha = 0,99$
0.90	0.96	0.99
0.95	0.96	0.99
1.00	0.95	0.99
1.50	0.89	0.96
2.00	0.83	0.93
2.50	0.78	0.90
3.00	0.74	0.86
3.50	0.71	0.83
4.00	0.67	0.80

Fonte: Pessoa e Nascimento Silva, 1998.

Sendo assim, percebe-se que para um EPA ampliado igual a 2, um IC de $\alpha=0,05$ ignorando o plano amostral, corresponde a um IC de $\alpha=0,17$ (1-0,83) caso fosse utilizado o plano amostral verdadeiro. Concluindo assim que para:

- EPAs de valor unitário, o nível de significância verdadeiro é igual ao nominal.
- EPAs maiores que 1, o nível de significância verdadeiro é superior ao nominal.
- EPAs menores que 1, o nível de significância verdadeiro é inferior ao nominal.

5.2.1.1.

Métodos para estimação dos parâmetros

A estimação dos parâmetros do modelo de regressão pode ser feita através do Método de Máxima Verossimilhança (MMV). Porém, dessa forma, o desenho amostral e os pesos não são considerados.

No caso do MMV, a função de verossimilhança amostral pode ser expressa sob a forma:

$$L(b) = f(y_1, \dots, y_n; b) = \prod_{i=1}^n f(y_i, b), \text{ onde;}$$

$y_i = (y_{i1}, y_{i2}, y_{i3}, \dots, y_{in})$ é um vetor $n \times 1$ formado por variáveis aleatórias independentes e identicamente distribuídas (IID), com função de densidade de probabilidade $f(y; \beta)$, onde β é o vetor $k \times 1$ de parâmetros desconhecidos a serem estimados

A partir desta função de verossimilhança, pode-se encontrar o estimador de MV de β , através da sua maximização. Matematicamente pode-se maximizar o logaritmo natural da função de verossimilhança, conhecida por log-verossimilhança e expresso da seguinte forma:

$$l(b) = \sum_{i=1}^n \log[f(y_i; b)]$$

Igualando as derivadas parciais de $l(\beta)$ com relação a cada componente de β a zero, obtém-se a maximização desta função:

$$\frac{\partial l(b)}{\partial b} = \sum_{i=1}^n u_i(b) = 0,$$

$$\text{onde, } u_i(b) = \frac{\partial \log[f(y_i; b)]}{\partial b} \text{ é a função de escores.}$$

A solução deste sistema determina as estimativas que compõe o vetor do estimador $\hat{b}$, estimador de máxima verossimilhança de β . Para uma discussão mais detalhada do MMV para ajustes de modelos paramétricos regulares, veja por exemplo, Garthwarte, Jolife e Jones (1995).

Porém, como descrito acima, esta metodologia MMV não considera a complexidade do plano amostral. A metodologia que leva em conta este argumento, e que é adotada pelo pacote SUDAANN, **é o método de Máxima Pseudo-Verossimilhança (MPV)**⁷. Por isso, uma sucinta descrição do método é apresentada a seguir.

⁷ Seria possível chegar as mesmas estimativas calculando os estimadores via Mínimos Quadrados Ponderados (MQP).

➤ Considere-se $T(b) = \sum_{i \in u} u_i(b)$ a soma dos escores, constituindo-se

em um vetor de totais populacionais. Para estimar este vetor de totais, pode-se usar um estimador linear ponderado da seguinte forma:

$$\hat{T}(b) = \sum_{i \in s} w_i u_i(b),$$

onde,

w_i é o peso amostral da unidade i ;

$u_i(\beta)$ é a função de escores.

O estimador de MPV proposto por Binder (1983), pode ser obtido resolvendo a solução da equação abaixo para β :

$$\hat{T}(b_{MPV}) = \sum_{i \in s} w_i u_i(b) = 0.$$

Assim, os estimadores⁸ pontuais de MV e MPV para β e para o desvio padrão do termo aleatório (σ_e), sob a forma matricial, são respectivamente:

$$\hat{b} = (x_s' x)^{-1} x_s' y_s \quad \text{e} \quad s_e = n^{-1} (y_s - x_s \hat{b})' (y_s - x_s \hat{b})$$

$$\hat{b} = (x_s' w_s x)^{-1} x_s' w_s y_s \quad \text{e} \quad s_e^w = (1_s' w_s 1_s)^{-1} (y_s - x_s \hat{b}_w)' w_s (y_s - x_s \hat{b}_w)$$

Onde, $w_s = \text{diag}[w_{i1}, w_{i2}, \dots, w_{in}]$, é a diagonal da matriz $n \times n$ com os pesos dos elementos da amostra na diagonal principal e 1_s é o vetor de dimensão $n \times 1$ de valores unitários. Para obtenção das variâncias estimadas se usa o artifício da Linearização de Taylor, adotado pelo SUDAAN.

No próximo passo da dissertação o interesse é estimar um modelo logístico. A equação de verossimilhança populacional é então dada por:

$$\sum_{i \in U} [y_i - P(x_i' b)] x_i = 0,$$

onde x_i é o vetor de variáveis explicativas da i -ésima observação e

$$P(x_i' \beta) = P(Y_i = 1 | X_i = x_i) = \exp(x_i' \beta) / (1 + \exp(x_i' \beta)).$$

⁸ Maior informação para o cálculo dos estimadores pode ser observada em Nascimento Silva, 1996.

A equação de pseudo verossimilhança, da qual são estimados os coeficientes β , é a seguinte:

$$\sum_{i \in S} w_i (y_i - p(x_i' \beta)) x_i = 0, \text{ onde } w_i \text{ é peso da } i\text{ésima observação.}$$

Uma vez explicados os métodos subjacentes ao modelo logístico e aos procedimentos do SUDAAN, os resultados são, então, apresentados e analisados.

5.2.2.

O modelo logístico

As variáveis que serão estudadas no modelo econométrico para identificação de fatores que têm relação com a probabilidade de uma pessoa ter plano de saúde serão: sexo (SEXO), nível educacional (EDUCA), grupo de idade (GIDADE), auto-avaliação do estado de saúde (SAUDE), região onde vive o indivíduo (REG), a presença de doença crônica (DCRONICA), faixa de renda familiar per capita (RENDA). A variável dependente é (PLANO), dicotômica, que assume valor 1 se o indivíduo tem plano e 0, em caso contrário.

Com a observação das relações existentes entre as variáveis explicativas e a variável dependente, espera-se poder conhecer os fatores que mais explicam as chances de um indivíduo possuir plano de saúde. Assim, os níveis dos fatores que se mostram mais favoráveis às chances de uma pessoa ter plano de saúde pode indicar a ocorrência de seleção adversa. Por exemplo, se para o efeito idade chega-se a conclusão de que pessoas idosas têm maiores probabilidades de ter plano de saúde que pessoas jovens, este resultado, poderá mostrar a presença de seleção adversa. Ou, ao menos, pode indicar que as seguradoras precisam construir defesas contra esta pressão desfavorável, visto que pessoas idosas, teoricamente, gerariam gastos maiores para as seguradoras do que indivíduos jovens.

Foram consideradas, também, interações entre as variáveis sexo e auto-avaliação do estado de saúde (SEXO*SAUDE), sexo e grupo etário (SEXO*GIDADE), nível educacional com doença crônica (EDUCA*DCRONICA) e nível educacional com percepção do estado de saúde

(EDUCA*SAUDE). O principal motivo de inclusão destas interações no modelo é que são combinações cujos efeitos podem influenciar as probabilidades dos indivíduos adquirirem planos de saúde. Por exemplo, pode-se esperar que pessoas com alto grau de instrução tenham chances diferenciadas de ter plano de saúde em função de terem ou não doença crônica, ou que pessoas do sexo feminino tenham chances diferentes de ter plano para grupos etários distintos. Na análise dos modelos estimados serão feitos outros comentários sobre as interações.

Na tabela 5, são apresentados os níveis das variáveis estudadas. O nível de comparação para todas as variáveis é o último nível, com exceção da variável DCRONICA, cujo nível de comparação é o primeiro (Têm doença crônica).

Tabela 5 – Variáveis usadas na regressão e suas categorias

Sexo	Homem = 1
	Mulher = 2
Educa	Sem escolaridade = 1
	1 a 8 anos = 2
	9 a 14 anos = 3
	15 anos ou mais = 4
Gidade	0 a 14 anos = 1
	15 a 50 anos = 2
	51 ou mais = 3
Saúde	Muito bom ou bom = 1
	Regular = 2
	Ruim ou muito ruim = 3
REG	Norte = 1
	Nordeste = 2
	Sudeste = 3
	Sul = 4
	Centro-oeste = 4
Dcronica	Têm = 1
	Não têm = 2
Renda	Até R\$ 100,00 = 1
	R\$ 100,00 a R\$ 200,00 = 2
	R\$ 200,00 a R\$ 500,00 = 3
	Acima de R\$ 500,00 = 4

Estas variáveis foram selecionadas em decorrência de se ter verificado, na literatura sobre seleção adversa, que as mesmas são as mais importantes para explicação do fenômeno.

Alternativamente, poderia-se ter buscado tais variáveis pela realização de regressões passo-a-passo porém, como este tipo de metodologia para obtenção do modelo ótimo deixa de considerar conceitos teóricos sobre o assunto, optou-se por não utilizá-lo aqui.

A regressão logística foi escolhida por ser apropriada para o tipo de situação onde a variável de resposta é binária. Portanto, será estimado um modelo logístico por intermédio do pacote estatístico SUDAAN, que permite a incorporação do plano amostral da PNAD 98.

Da forma proposta, a probabilidade de uma pessoa ter plano de saúde pode ser modelada da seguinte maneira:

$$P_i = E(Y_i=1 | X_{1i}, X_{2i}, \dots, X_{ni}) = 1 / (1 + e^{-(X_i\beta)}) \quad i = 1, 2, \dots, n$$

onde X_{ji} 's são os regressores dos i -ésimo indivíduo

$Y_i = 1$, pessoa com plano;

0, caso contrário.

Esta expressão é equivalente a: $P_i = 1 / (1 + e^{-Z})$, função de distribuição logística, onde $Z = X\beta$ e X é o vetor de variáveis explicativas e β o vetor de parâmetros.

Se a probabilidade de uma pessoa possuir plano é P_i , então a probabilidade de não possuir é $(1-P_i)$ e $(1-P_i) = 1 / (1 + e^Z)$, de modo que é possível escrever que $P_i / (1-P_i) = (1 + e^Z) / (1 + e^{-Z})$. Assim, $P_i / (1-P_i)$ é simplesmente a razão de probabilidade (em inglês, *odds ratio*) em favor de um indivíduo ter plano de saúde. É a razão entre a probabilidade de um indivíduo ter plano e a probabilidade dele não ter plano.

O log da razão de probabilidades L_i ,.

$$L_i = \ln \left(\frac{P_i}{1-P_i} \right) = Z_i = X_i\beta \quad \text{é o logit } L, \text{ daí o nome modelo logit}$$

O modelo inicial de regressão logística, a partir do qual a análise será baseada tem a seguinte equação:

Modelo1:

$$\ln\left(\frac{P_{ijklmn}}{1-P_{ijklmn}}\right) = m + b_i \text{sexo} + b_j \text{educa} + b_k \text{gidade} + b_l \text{saude} + b_m \text{reg} + \\ b_n \text{dcronica} + b_o \text{renda} + b_{kn} \text{dcronica} * \text{gidade} + b_{il} \text{sexo} * \text{saude} + \\ b_{jn} \text{dcronica} * \text{educa} + b_{jl} \text{saude} * \text{educa} + b_{jl} \text{renda} * \text{educa} + x_i$$

onde,

b's são os coeficientes das variáveis e:

i=1 se sexo = homem e *i*=2 se sexo = mulher

j=1 se nível de educação = s/ instrução; *j*=2 se educação = 1 a 8 anos de estudo; *j*=3 se nível de educação = 9 a 14 anos de estudo; *j*=4 se nível de educação é 15 anos ou mais.

k=1 se idade = 0 a 14 anos; *k*=2 se idade = 15 a 50 anos; *k*=3 se idade é acima de 50 anos

l=1 se saúde = muito boa ou boa; *l*=2 se saúde = regular; *l*=3 se saúde=ruim

m=1 se região = norte; *m*=2 se região = nordeste; *m*=3 se região = sudeste; *m*=4 se região = sul; *m*=5 se região = centro-oeste.

n=1 se não tem doença crônica; *n*=2 se tem doença crônica

o=1 se renda = até R\$ 100,00; *o*=2 se renda = de R\$ 100 a R\$ 200,00;

o=3 se renda = de R\$ 200 a R\$ 500,00; *o*=4 se renda = acima de R\$ 500,00

o modelo será estimado, como já ressaltado, com auxílio da PROC RLOGIST, do pacote SUDAAN, usada em ambiente SAS.

5.2.3.**Análise da regressão logística**

Os testes de Wald e os respectivos p-valores fornecido pelo SUDAAN para os coeficientes dos efeitos incluídos neste modelo estão na tabela 6:

Tabela 6 – Estatísticas de Wald para o modelo 1.

Contrast	Degrees of		P-value	
	Freedom	Wald F	Wald F	
OVERALL MODEL	32	433.07	0.0000	
REG	4	74.49	0.0000	
SEXO * SAUDE	2	0.44	0.6410	
SEXO * GIDADE	2	72.43	0.0000	
DCRONICA * EDUCA	3	18.79	0.0000	
SAUDE * EDUCA	6	6.60	0.0000	
RENDIA * EDUCA	9	119.76	0.0000	

Fonte: PNAD/1998.

Com a tabela 6 observa-se que o teste F de Wald apresentou valor significativo para todas as interações, com exceção da interação SEXO*SAUDE, que apresentou valor 0,44 para o teste de Wald e P-valor associado de 0,6410, não significativo para os níveis usuais de análises estatísticas. Os valores estimados das razões de chance (*odds ratios*), bem como dos coeficientes do modelo 1 encontram-se no anexo.3.

O que se conclui, então, é que não há mudanças nas probabilidades ou nas chances de pessoas com percepções diferentes do próprio estado de saúde (auto-avaliação do estado de saúde) em função dessas pessoas serem do sexo masculino ou feminino.

Logo, em virtude da obtenção deste valor não significativo para a interação SEXO*SAUDE, foi estimado um outro modelo, doravante denominado de modelo 2, e que se distingue do primeiro modelo estimado pela ausência da interação SEXO*SAUDE.

Este segundo modelo, como pode ser visto na tabela 7, apresentou todos os efeitos e interações significativos ao nível de significância de 0,01.

Tabela 7 - Estatísticas de Wald para o modelo 2

Contrast	Degrees of		P-value	
	Freedom	Wald F	Wald F	
OVERALL MODEL	37	455.10	0.0000	
REG	4	74.50	0.0000	
SEXO * GIDADE	2	77.52	0.0000	
DCRONICA * EDUCA	3	18.81	0.0000	
SAUDE * EDUCA	6	6.58	0.0000	
RENDIA * EDUCA	9	119.79	0.0000	

Fonte: PNAD/1998.

Para este segundo modelo todos os efeitos foram significativos ao nível de significância de 0,01. O teste de Wald para o modelo global também foi significativo. Todos os p-valores foram menores que 0,0001. Com a validação dos coeficientes, há condições de se avaliar o modelo e as relações entre as variáveis e o significado de cada efeito simples, bem como o das interações. As tabelas 8 e 9, adiante, contêm os valores dos coeficientes estimados para os efeitos principais e interações do modelo 2, respectivamente.

Tabela 8 – Coeficientes estimados para o modelo 2

Efeito	Nível	Nome	Coef. Beta
Intercepto		Intercepto	-0,04
Sexo	1	Homem	-0,35
	2	Mulher*	0,00
Dcrônica	1	Não Tem	0,27
	2	Tem*	0,00
Saúde	1	Muito boa/boa	0,84
	2	Regular	0,66
	3	Ruim*	0,00
Gidade	1	0 a 14 anos	0,38
	2	15 a 50 anos	-0,21
	3	acima 50 anos*	0,00
Educa	1	S/ instrução	-3,77
	2	1 a 8 anos de estudo	-2,85
	3	9 a 14 anos de estudo	-1,32
	4	15 ou mais anos estudo*	0,00
Reg	1	Norte	-0,06
	2	Nordeste	-0,12
	3	Sudeste	0,45
	4	Sul	0,07
	5	Centro Oeste*	0,00
Renda	1	Até R\$ 100,00*	0,00
	2	R\$ 100 a R\$ 200	-1,04
	3	R\$ 200 a R\$ 500	-0,14
	4	Acima de R\$ 500	-0,92

Fonte: PNAD/1998.

*Níveis de referência

Tabela 9 – Coeficientes das interações do modelo 2

Efeito	Coef. Beta
sexo 1 * Gidade 1	0,37
sexo 1 * Gidade 2	0,12
Dcrônica 1 * Educa 1	-0,35
Dcrônica 1 * Educa 2	-0,05
Dcrônica 1 * Educa 3	-0,03
Saúde 1 * Educa 1	-0,30
Saude 1 * Educa 2	-0,68
Saude 1 * Educa 3	-0,63
Saude 2 * Educa 1	-0,37
Saude 2 * Educa 2	-0,54
Saude 2 * Educa 3	-0,58
Renda 2 * Educa 1	2,54
Renda 2 * Educa 2	2,13
Renda 2 * Educa 3	1,37
Renda 3 * Educa 1	2,71
Renda 3 * Educa 2	2,08
Renda 3 * Educa 3	1,29

Fonte: PNAD/1998.

A partir destes resultados, algumas observações interessantes e importantes relacionadas ao problema da seleção adversa podem ser obtidas. É possível, por exemplo, verificar a influência das diversas covariáveis estudadas na probabilidade de um indivíduo ter plano de saúde.

Vale ressaltar que, dado que algumas interações foram significativas, a análise de cada variável deve ser feita considerando-se os efeitos das demais variáveis que com ela interagem.

Para avaliação dos efeitos da variável renda, é preciso considerar que a mesma interage com o nível educacional (variável EDUCA). Esta, por sua vez, causa impactos na probabilidade de uma pessoa ter plano de saúde de forma diferenciada, dependente do fato do indivíduo ter ou não ter doença crônica e também da auto-avaliação do estado de saúde.

Assim, foram calculadas as probabilidades de se ter cobertura por plano de saúde para indivíduos sem doenças crônicas e com boa auto-avaliação do estado de saúde. Os resultados estão na tabela 10:

Tabela 10- Chances de ter plano contra não ter por níveis de renda e educação para pessoas sem doenças crônicas e boa auto avaliação do estado de saúde.

Níveis de renda e educação	Chances de ter plano contra não ter
P(renda até 100, educa = 1)	0,07
P(renda de 100 a 200, educa = 1)	0,30
P(renda de 200 a 500, educa = 1)	0,88
P(renda acima de 500, educa = 1)	0,17
P(renda até 100, educa = 2)	0,17
P(renda de 100 a 200, educa = 2)	0,50
P(renda de 200 a 500, educa = 2)	1,17
P(renda maior que 500, educa = 2)	0,42
P(renda até 100, educa = 3)	0,78
P(renda de 100 a 200, educa = 3)	1,08
P(renda de 200 a 500, educa = 3)	2,46
P(renda maior que 500, educa = 3)	1,95
P(renda até 100, educa = 4)	2,92
P(renda de 100 a 200, educa = 4)	1,03
P(renda de 200 a 500, educa = 4)	2,53
P(renda maior que 500, educa = 4)	7,32

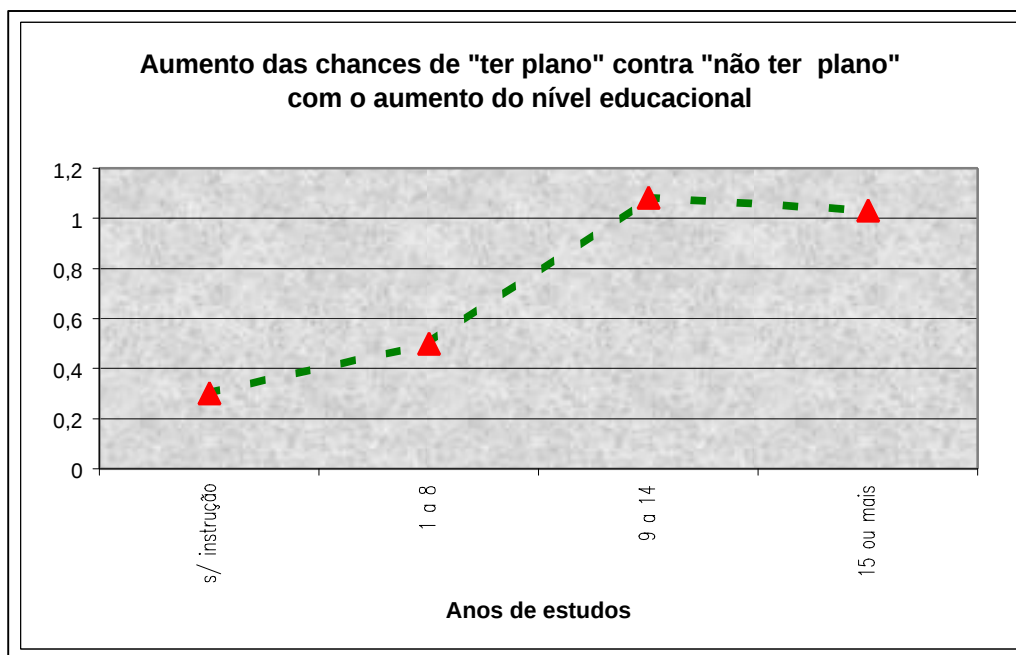
Fonte: PNAD/1998.

Com a análise da tabela 10 é possível perceber que para pessoas sem doença crônica e com boa auto-avaliação do estado de saúde a probabilidade de ter plano de saúde é muito maior para aqueles com maiores níveis de instrução e renda. Indivíduos com renda familiar per capita igual ou superior a R\$ 500,00 têm cerca de 104 vezes ($=7,32/0,07$) mais probabilidades de ter cobertura de plano de saúde do que indivíduos com renda familiar per capita de até R\$ 100,00 e sem instrução. Para estas pessoas sem instrução, as chances de ter plano de saúde são sempre menores que as de não ter, independentemente do nível de renda.

Quando a análise foca os que têm rendimentos familiar per capita entre 100 e 200 reais nota-se que as probabilidades de possuir plano são mais elevadas para os que têm mais anos de estudos. Para pessoas nesta faixa de rendimento, os que estudaram por mais de 15 anos têm 2 vezes mais chances de possuir cobertura de plano de saúde do que os indivíduos que declararam ter de 1 a 8 anos de estudos. Se a comparação é feita com o grupo de indivíduos sem instrução, as

chances de ter plano são ainda maiores para os que têm mais instrução, chegando esta diferença a aproximadamente 3,5 vezes. O gráfico abaixo ilustra as chances maiores de plano para os que têm mais anos de estudos.

Figura 5 - Aumento da chance de ter plano contra não ter, em função do aumento do nível educacional dado o rendimento familiar per capita de 100 a 200 reais.



Fonte: PNAD/1998

O nível educacional foi incluído no modelo porque se considerou que pessoas com maior nível de instrução talvez possam perceber a necessidade ou a importância de ter cobertura de plano de saúde de forma diferente de pessoas com baixo grau de instrução.

Como já mencionado, a variável EDUCA foi incorporada ao modelo interagindo com as variáveis SAÚDE, DCRONICA e RENDA. A interação de EDUCA com SAÚDE teve por objetivo verificar se pessoas com elevado nível de instrução e identificação do próprio estado de saúde como sendo “ruim” ou “muito ruim” têm chances maiores de ter plano de saúde. Contudo, para avaliação deste resultado é necessário que os níveis da variável DCRONICA e RENDA, que também interagem com a variável EDUCA, sejam mantidos fixos.

Com este procedimento, é possível perceber como variam as chances de indivíduos com doenças crônicas e com auto avaliação ruim do estado de saúde terem plano de saúde para combinações de renda e nível educacional.

As razões de vantagens (*odd ratios*) para estas comparações podem ser conhecidas dividindo-se as chances da tabela 11.

Tabela 11 – Chances de ter cobertura de plano contra não ter, para pessoas com doenças crônicas e percepção ruim do estado de saúde.

Combinações dos níveis de renda e educação	Chances de ter plano contra não ter
P(renda até 100, educa = 1)	0,02
P(renda de 100 a 200, educa = 1)	0,10
P(renda de 200 a 500, educa = 1)	0,29
P(renda acima de 500, educa = 1)	0,06
P(renda até 100, educa = 2)	0,06
P(renda de 100 a 200, educa = 2)	0,17
P(renda de 200 a 500, educa = 2)	0,39
P(renda maior que 500, educa = 2)	0,14
P(renda até 100, educa = 3)	0,26
P(renda de 100 a 200, educa = 3)	0,36
P(renda de 200 a 500, educa = 3)	0,81
P(renda maior que 500, educa = 3)	0,64
P(renda até 100, educa = 4)	0,96
P(renda de 100 a 200, educa = 4)	0,34
P(renda de 200 a 500, educa = 4)	0,84
P(renda maior que 500, educa = 4)	2,41

Fonte: PNAD/1998.

Com os *odds ratios* calculados a partir da tabela 11, percebe-se que as pessoas com maiores rendimentos têm maiores chances de ter plano de saúde. Logo, mesmo para grupo de indivíduos com doenças crônicas e auto avaliação “ruim” do estado de saúde o comportamento é semelhante ao dos indivíduos com boa auto avaliação do estado de saúde. Dividindo-se o valor de 2,41, correspondente a chance de ter plano contra não ter para pessoas com renda acima de 500 reais e nível mais alto de educação, por 0,14 (renda maior que 500 reais e nível 2 de educação) verifica-se que as chances em favor do maior nível educacional, no caso, em favor dos indivíduos com mais de 15 anos de estudos,

são 17 vezes maiores. Ou seja, quanto maior o nível educacional, maior a chance de o indivíduo ter plano de saúde. É válido enfatizar que, neste caso, o efeito renda está eliminado, já que a comparação foi feita para indivíduos com a mesma faixa de rendimentos.

Com relação às variáveis explicativas do efeito sexo e idade, é possível inferir sobre as razões de chances obtidas quando ocorre mudança no nível da variável sexo (SEXO) para níveis constantes da variável de faixa etária (IDADE). As razões de chances (*odds ratios*) calculados para estes grupos facilitam a análise.

Estas razões de chances foram calculadas comparando-se as chances de cobertura por plano para os grupos estudados, mostrados na tabela 12, abaixo:

Tabela 12 – Estimativas das chances de ter plano contra não ter.

Combinações de sexo e idade	Chances de ter plano contra não ter
P(homem idade ate 14 anos)	2,36
P(mulher idade ate 14 anos)	1,40
P(homem 15<= idade<=50)	0,81
P(mulher 15<= idade<=50)	0,78
P(homem idade>50)	0,68
P(mulher idade>50)	0,96

Fonte: PNAD/1998.

Assim, quando o foco de investigação passa a ser a faixa etária e sua relação com a variável sexo para explicação da probabilidade de ter plano, observa-se que mulheres têm maiores probabilidades de estarem cobertas por planos do que homens para o grupo etário acima de 50 anos. A tabela 12, com os *odds ratios* calculados a partir da tabela 11, estão a seguir:

Tabela 13 – Estimativas das razões de vantagens quando são variados os níveis de sexo para níveis fixos de idade.

Idade	Mudança de Nível de Sexo	Razões de Chance
Até 14 anos	H para M	0,595
De 15 a 50 anos	H para M	0,961
Acima de 50 anos	H para M	1,419

Fonte: PNAD/1998.

É interessante observar que as mulheres chegam a ter 40% mais chances de ter plano de saúde do que homens quando se define o grupo estudado como sendo o de pessoas com mais de cinquenta anos de idade. Este resultado pode indicar uma maior preocupação com a saúde de mulheres dessa faixa etária.

Quando a análise é feita em separado para homens e mulheres nota-se que as maiores probabilidades de cobertura são para crianças até 14 anos. Para as mulheres desta faixa etária, as chances de ter plano de saúde são 45% maiores que as chances de mulheres com idade acima de 50 anos. Para os homens de até 14 anos, as chances de ter plano, se comparadas às dos homens com mais de 50 anos, são 3,5 vezes maiores.

Com a estimação do modelo logístico pode-se perceber que os idosos não constituem o grupo de indivíduos com maiores chances de estarem cobertos por plano.

Pessoas com doenças crônicas têm ligeiramente mais probabilidades de contratarem planos de saúde que pessoas sem problemas crônicas.

Tabela 14 - Razões de chances para pessoas sem doença crônica e com doença crônica, com 9 a 14 anos de estudos e boa auto avaliação do estado de saúde para diferentes níveis de renda.

Razões de chances para pessoas sem doença crônica e com doença crônica.	Razões de chances
Pessoas s/ doença crônica Renda=1 saúde=1 educação=3	0,40
Pessoas s/ doença crônica Renda=2 saúde=1 educação=3	0,56
Pessoas s/ doença crônica Renda=3 saúde=1 educação=3	1,27
Pessoas s/ doença crônica Renda=4 saúde=1 educação=3	0,77
Pessoas c/ doença crônica Renda=1 saúde=1 educação=3	0,31
Pessoas c/ doença crônica Renda=2 saúde=1 educação=3	0,43
Pessoas c/ doença crônica Renda=3 saúde=1 educação=3	0,97
Pessoas c/ doença crônica Renda=4 saúde=1 educação=3	0,59

A análise da tabela 14 acima indica que pessoas sem doenças crônicas têm maiores probabilidades de terem planos de saúde que pessoas com doenças crônicas para todos os níveis de renda. A título de ilustração, por exemplo, dividindo-se 0,40 (odds ratio para renda=1, sem doença crônica) por 0,31 (odds ratio para renda =1, com doença crônica), encontra-se 1,29, o que indica que as pessoas do grupo sem doenças crônicas têm 29% a mais de chances de ter plano de saúde que pessoas com doenças crônicas. Estas chances maiores favoráveis às pessoas em bom estado de saúde se repetem para todos os níveis de renda.

Esta análise, somada ao estudo descritivo de dados da PNAD, mostra que, apesar de existirem “pressões” adversas para as administradoras de planos de saúde, as precauções e medidas tomadas para não deixar a seleção adversa afetá-las, como por exemplo a diferenciação de preços por faixa etária, pode estar surtindo efeitos favoráveis.

A partir da seção seguinte, o foco passa a ser o risco moral.

5.3. Análise do risco moral

Nas seções anteriores foi explicitado o significado do risco moral. Além disso, foram feitos comentários sobre a importância de se considerar este fenômeno para o entendimento da dinâmica de funcionamento de diversos cenários econômicos e também se realizou uma análise descritiva sobre o mesmo na área de saúde, a partir do banco de dados da PNAD 98.

Agora, nesta seção, vamos procurar estudar mais objetivamente o risco moral, a fim de que possamos encontrar evidências estatísticas que possam vir a contribuir para o entendimento do assunto e complementar as suspeitas encontradas até aqui. Para isso é proposta a construção do indicador de risco moral - o IRM. O mecanismo de construção deste indicador é descrito em detalhes mais a frente.

Como já foi mencionado anteriormente, é fundamental, para a correta análise do IRM que, na construção do mesmo, sejam consideradas as características do plano amostral da PNAD 98.

5.3.1. Construção do Indicador de Risco Moral - IRM

Com o objetivo de identificar indícios da presença de risco moral em um determinado mercado propõe-se construir um indicador de risco moral, doravante denominado IRM. Espera-se que o IRM seja um indicador dicotômico, construído a partir um teste de hipótese estatístico, e que venha a fornecer, fundamentando-se na teoria de estimação estatística, uma indicação sobre a existência ou não do risco moral, levando em conta a complexidade do plano amostral da PNAD 98.

Como exemplo de construção, e interpretação do IRM aplicar-se-á a metodologia proposta ao mercado de planos de saúde brasileiro.

Destaca-se ainda que, também com base em informações do Suplemento Saúde da PNAD 98, modelos econométricos de regressão serão usados para permitir uma espécie de confirmação do risco moral. A vantagem da construção do indicador IRM é que este é de simples entendimento e menos complexo que a

abordagem econométrica tradicional, pois apenas algumas poucas variáveis são usadas para sua obtenção.

Para construção do IRM, serão estimados intervalos de 95% de confiança para o número de consultas realizadas por pessoas com planos (NCP) e para o número de consultas realizadas por pessoas sem cobertura de planos de saúde (NCSP) no último ano. Será incorporando na análise o desenho amostral da PNAD 98 que, como já foi dito, elimina o vício das estimativas e já teve sua importância citada.

Para os dois grupos estudados, não foram consideradas pessoas com doenças crônicas. Tal decisão se fundamenta no fato de que pessoas com doenças crônicas realizam visitas médicas ou fazem consultas periódicas em decorrência da necessidade de cada indivíduo com relação ao problema de saúde. Evidentemente, esta situação não pode ser considerada como importante na análise do risco moral.

Já para o grupo de pessoas sem cobertura de planos de saúde optou-se por considerar as que declararam rendimento⁹ mensal de todas as fontes como sendo superior a seis mil reais. A justificativa para esta barreira na variável ‘renda’ deve-se a busca da eliminação do efeito da desigualdade social na análise do risco moral. Ou melhor, tenta-se identificar o comportamento com relação a visitas médicas de pessoas cujo fator ‘renda’ não seja um impedimento para a realização de consultas, de forma que caso seja observado um menor número de consultas por parte deste grupo de indivíduos com relação às pessoas que têm plano de saúde, seja possível dizer que a maior taxa de realização de consultas médicas pelas pessoas com plano seja decorrente somente do risco moral.

⁹ Para a construção da variável rendimento foi considerada renda de todas as fontes.

5.3.2. A interpretação do IRM

Com o indicador de risco moral – IRM, espera-se dizer, com um determinado nível de confiança estatística, se o risco moral está ou não presente no mercado de planos de saúde.

A princípio, o risco moral na área de saúde se caracteriza e fica bem determinado pelo maior número de consultas médicas realizadas por pessoas cobertas por seguro saúde ou plano de saúde. Logo, espera-se que pessoas com planos de saúde realizem, *ceteri paribus*¹⁰, num determinado período de tempo, um número maior de consultas do que pessoas sem plano de saúde ou não cobertas por qualquer seguro.

No início do parágrafo anterior o termo “a princípio” foi usado para permitir a lembrança de que algumas considerações devem ser feitas conjuntamente com a suposição contida no mesmo. Por exemplo, pode ser que pessoas não cobertas por planos de saúde realizem menos consultas por motivos de restrições orçamentárias. Da mesma forma, também pode ser possível que pessoas que têm planos de saúde sejam pessoas mais preocupadas com a saúde do que as que não têm planos e talvez realizem mais consultas médicas em função desta preocupação. Evidentemente todos esses pontos devem ser levados em consideração para um parecer definitivo e análise dos resultados.

Como o IRM se fundamenta na comparação do número de consultas realizadas por pessoas cobertas e pessoas não cobertas por planos de saúde, não são considerados os fatores citados acima, referentes às diferentes situações que possam interferir no número de consultas dos indivíduos, exceto pela variável “ter plano”. E é por isso que o IRM é apenas considerado como sendo um “*indicador*” da existência de risco moral.

Contudo, devido à importância e a praticidade de cálculo do IRM, este pode ser visto como de grande utilidade para a detecção do risco moral. No caso de necessidade de estudos de diversos mercados simultaneamente, como diferentes cidades, por exemplo, o IRM pode funcionar como um filtro, podendo ser largamente utilizado e permitindo uma indicação a priori dos mercados onde a

¹⁰ Assume-se, neste momento, que todas as outras variáveis que possam vir a influir no número de consultas médicas (sexo, idade, renda, etc) estejam igualmente distribuídas entre os dois grupos.

indicação do risco moral é positiva. Apenas para exemplificar um ganho com o uso do IRM, com a filtragem preliminar, é possível se deter aos cenários que apresentam IRM indicando risco moral e assim estudos mais aprofundados, inclusive com o complemento da abordagem econométrica, podem ser realizados apenas nestes mercados onde a indicação do risco moral através do IRM se verifica.

Além das comparações entre diferentes subpopulações e comparações regionais o IRM também pode ser aplicado a outros mercados tais como os de seguros de automóveis, residências e todos aqueles em que a teoria econômica vislumbra a possibilidade de ocorrência do fenômeno.

Basicamente, a idéia do IRM é comparar os dois intervalos estimados, calculados levando-se em consideração o plano amostral. O risco moral é então avaliado com base na seguinte regra decisão:

- Se $NCSP < NCP$ então há indícios de que o risco moral está presente e IRM é positivo.
- Se $NCSP = NCP$ então não há indícios estatísticos da presença de risco moral.

Onde,

NCP é o número médio de consultas de pessoas com plano de saúde e

$NCSP$ é o número médio de consultas de pessoas sem plano de saúde.

Ou seja, não havendo interseção dos intervalos estimados, podemos dizer que o risco moral ocorre.

5.3.3.**Estimação do IRM para o mercado brasileiro de planos de saúde**

A partir dos dados da PNAD 98, foi possível a construção de um teste de hipóteses estatístico para a avaliação do IRM, cuja formulação é a seguinte:

$$\begin{cases} H_0 : E(NCP) = E(NCSP) \\ H_1 : E(NCP) > E(NCSP) \end{cases}$$

onde,

$E(NCP)$ é o número médio de consultas de pessoas com plano de saúde e

$E(NCSP)$ é o número médio de consultas de pessoas sem plano de saúde.

Para o teste acima, se admite a hipótese da normalidade para a distribuição da média amostral do número de consultas e um nível de significância α de 0,05.

Com a realização do teste de hipóteses especificado acima como o auxílio do pacote SUDAAN chega-se as seguintes estatísticas:

Para o Número de consultas de pessoas com plano no nível Brasil (NCP):

$$E(NCP_{\text{Brasil}}) = 2,30$$

$$\text{Desvio Padrão} = 0,02$$

A partir deste valor médio e seu desvio padrão obtém-se o intervalo de confiança desejado, mostrado abaixo:

$$\text{Intervalo estimado para o } NCP_{\text{Brasil}} = [2,26 ; 2,33]$$

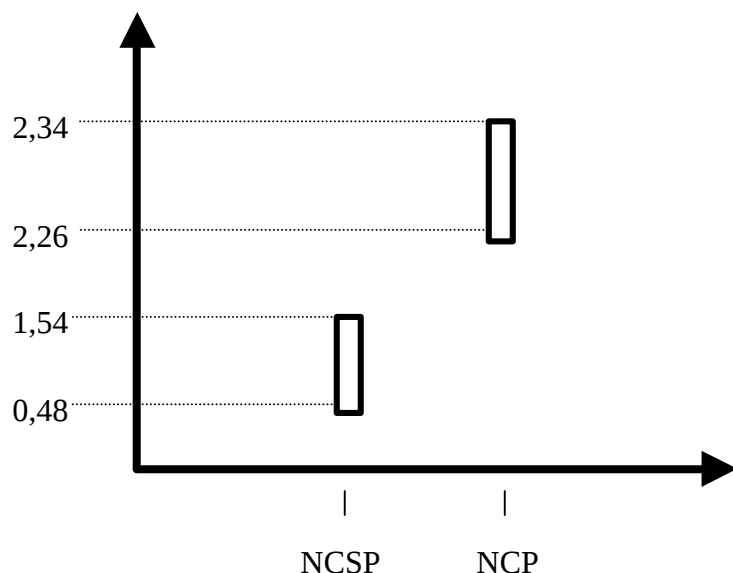
Para o Número de consultas de pessoas sem plano no nível Brasil e rendimento superior a seis mil reais (NCSP):

$$E(NCSP)_{\text{Brasil}} = 1,01$$

$$\text{Desvio Padrão} = 0,27$$

$$\text{Intervalo estimado para o } NCSP_{\text{Brasil}} = [0,48 ; 1,53]$$

Figura 6 – Intervalo de confiança para o número de consultas de pessoas com plano e sem plano de saúde.



Com estes resultados, nota-se que para o nível Brasil, não há interseção entre os intervalos de confiança estimados, como ilustra a figura 75. Pode-se dizer que pessoas com plano de saúde realizam mais consultas médicas do que as que não têm. Como o grupo de indivíduos sem plano de saúde estudado é formado por pessoas que têm rendimento mensal acima de seis mil reais, pode-se supor que o menor número de consultas realizado por este grupo não é decorrente de restrições orçamentárias. Logo, o resultado positivo do IRM indica forte possibilidade de presença do risco moral no cenário brasileiro.

Além de análise no nível mais agregado nacional, também foram construídos intervalos de confiança para avaliação do risco moral nas regiões brasileiras. Apesar do estudo com maior abrangência (Brasil) ser interessante pela quantidade de dados disponíveis, o desejo de conhecimento do fenômeno no nível de regiões se baseia na possibilidade de análises posteriores com cruzamentos de variáveis mais facilmente obtidas e melhor caracterizadas para as regiões do país.

Então, foram construídos os indicadores IRM para cada uma das cinco regiões brasileiras utilizando-se da mesma metodologia de construção do IRM em nível nacional. Os resultados são apresentados nas tabelas seguintes:

Tabela 15 – Número médio de consultas de pessoas com plano(NCP)

Regiões	E(NCP)	Desvio padrão
Norte	1,90	0,07
Nordeste	2,47	0,06
Sudeste	2,30	0,30
Sul	2,28	0,05
Centro-Oeste	4,18	0,07

Fonte: PNAD/1998.

Tabela 16 – Número médio de consultas de pessoas sem plano e com rendimento superior a seis mil reais (NCSP).

Regiões	E(NCSP)	Desvio padrão
Norte	0,83	0,34
Nordeste	2,71	0,89
Sudeste	1,05	0,82
Sul	0,71	0,22
Centro-Oeste	0,76	0,26

Fonte: PNAD/1998.

Tabela 17 – Intervalos de confiança para NCP e NCSP por regiões.

	(NCP)	(NCSP)
Norte	[1,76 ; 2,04]	[0,16 ; 1,50]
Nordeste	[2,35 ; 2,59]	[0,97 ; 4,45]
Sudeste	[2,24 ; 2,36]	[0,00 ; 2,66]
Sul	[2,18 ; 2,38]	[0,28 ; 1,14]
Centro-Oeste	[2,04 ; 2,32]	[0,25 ; 1,27]

Fonte: PNAD/1998.

Pela observação das tabelas e dos respectivos valores estimados para números de consulta de pessoas com plano e de pessoas sem plano nas regiões do Brasil, verifica-se que nas regiões Norte, Sul e Centro-Oeste o IRM indica indícios da presença de risco moral.

Entretanto, para as regiões Nordeste e Sudeste, o número médio de consultas de pessoas sem plano de saúde não é estatisticamente diferente do número médio estimado de consultas de indivíduos com plano de saúde.

Este resultado é muito importante e parece contrariar a teoria sobre o risco moral na área de saúde. No mínimo, poderia dar sustentação a um debate mais detalhado sobre o tema. Esta discussão poderia ser pautada nas buscas a respostas para os seguintes questionamentos:

1 - O risco moral aparece na área de planos de saúde da mesma forma em que se apresenta em outros mercados como, por exemplo, o mercado de automóveis?

2 – Apesar de não haver dúvidas quanto aos mecanismos que explicam o risco moral, será que as pessoas que tem planos de saúde realizam consultas desnecessárias? Ou será que são as pessoas sem cobertura de planos que deixam de ir ao médico o quanto deveriam?

Este último questionamento surge porque, segundo informações veiculadas pelos conselhos regionais de medicina em todo o Brasil, as pessoas devem procurar auxílio médico em caso de suspeita de qualquer problema de saúde, por mais simples que estes possam parecer, de modo que, certamente, não é fácil a argumentação sobre a intensidade do risco moral na área de saúde.

Entretanto, para o estudo em nível regional, pelo fato das sub-populações de análise terem sido limitadas pelos valores de renda mensal superior a seis mil reais, não foram possíveis inferências mais confiáveis, em decorrência do pequeno tamanho amostral para indivíduos sem plano de saúde para esta faixa de rendimento.

Como a simples análise do IRM não permite uma argumentação mais conclusiva sobre o assunto, é proposta, num capítulo subsequente, a construção de um modelo econométrico “hurdle binomial negativo”. Espera-se, com este modelo, inferir sobre variáveis que interferem no número de consultas e, dessa forma, avaliar se é o fato de a pessoa ter ou não ter plano determinante das visitas médicas.

5.4.

Avaliação econométrica do risco moral

Antes de se iniciar os comentários sobre os modelos econométricos para avaliação do risco moral, considera-se importante a referência a alguns aspectos relacionados às técnicas estatísticas de análise e suas possíveis limitações em economia de saúde.

Por exemplo, além dos problemas relacionados à obtenção e adequabilidade dos dados disponíveis, fatores oriundos das limitações das técnicas estatísticas também devem ser considerados. As mais importantes limitações são as decorrentes de problemas referentes à endogeneidade e à causalidade. Estes dois conceitos serão comentados no tópico a seguir.

5.4.1.

Endogeneidade e causalidade

A endogeneidade é um problema que se observa quando uma (ou algumas) das variáveis explicativas é correlacionada ao erro aleatório. O problema de endogeneidade faz com que os modelos com esta característica não possam ser estimados por mínimos quadrados ordinários (MQO), visto que uma das hipóteses básicas do modelo clássico de regressão (MQO) pressupõe ausência de correlação entre uma variável independente (x_i) e o termo estocástico e_i , ou seja, $\text{Cov}[x_i, e_i] = 0$. Na presença de endogeneidade os estimadores de MQO são enviesados.

A endogeneidade pode ser originada pela ausência de alguma variável explicativa importante no modelo, o que carrega o termo aleatório e aumenta a probabilidade deste se correlacionar com as demais variáveis independentes.

Há, também, necessidade de se confirmar, a priori, se os parâmetros estimados do modelo refletem realmente uma relação causal entre as variáveis explicativas (x_i) e a variável resposta (y_i) e não uma correlação espúria, causada pela associação entre x_i e e_i . Ou seja, é preciso avaliar se a correlação entre x_i e y_i não é decorrente do fato de ambas estarem relacionadas a uma terceira variável, não presente na modelagem.

Em alguns casos, os problemas da endogeneidade e causalidade, associados a outros elementos limitadores, tais como a presença de variáveis dependentes discretas ou limitações para obtenção de dados, podem dificultar em muito a busca por soluções adequadas para os problemas em economia de saúde.

5.4.2.

Alternativas metodológicas para avaliação do risco moral

Apesar de não haver ainda a desejada unanimidade metodológica na estimação de modelos econométricos com micro-dados de saúde, avanços foram observados nos últimos anos neste campo.

Segundo Lisboa (2001), atribui-se grande parcela do desenvolvimento destas técnicas de estimação exatamente ao fato de não ser viável, ou ao menos recomendável, o uso de técnicas econométricas tradicionais na abordagem de economia de saúde. Uma das justificativas é que a demanda por serviços de saúde apresenta características que a diferenciam da demanda por outros tipos de bens. Ainda na opinião de Lisboa, duas características se destacam: a primeira diz respeito ao fato de que uma parcela considerável da população não apresenta demanda por serviços médicos ou gastos com serviços médicos. Neste caso há necessidade de se entender se a presença de zeros na amostra selecionada significa uma escolha individual ou está associada a um problema de seleção amostral, ou seja, a impossibilidade financeira de indivíduos consumirem serviços de saúde. Quando se trata, por exemplo, da estimação de modelos para conhecimento da demanda por serviços de saúde, a presença de zeros pode ser interpretada como resultado de uma escolha individual.

Uma outra questão relevante para Lisboa está relacionada à natureza da decisão. Ou seja, há necessidade de se verificar se o consumidor toma a decisão de consumir e do quanto consumir ao mesmo tempo. Por exemplo, em modelos de acesso normalmente a decisão de consumir é do paciente, mas a decisão sobre o quanto utilizar o serviço médico ou o quanto consumir de exames é do médico.

Em modelos de determinação de gastos com serviços de saúde, ou modelos de estimação de demanda de serviços de saúde, nem sempre é possível identificar os dois processos, existindo um debate na literatura sobre qual o

modelo mais indicado em cada caso. Verifica-se na literatura que muitos autores tratam este tipo de situação com a utilização de algumas técnicas de correções ou ajustes que consideram a existência de tais características, tais como a correção de Heckman, sugerida por Heckman (1979), os modelos tobit de dois estágios ou com barreira (*hurdle*), difundidos em muitos artigos. Essas técnicas variam em função do tipo de abordagem empregada. A técnica de Heckman, por exemplo, é útil para corrigir viés de seleção. Entretanto, a correção de Heckman não se aplica a modelos de dados de contagem

Como um desses exemplos podemos citar o artigo de Duan *et al* (1983), onde é feita uma comparação de modelos alternativos para demanda por serviços médicos. Dos comentários encontrados neste artigo se extrai que os modelos com barreira (*hurdle*) podem ser empregados para estimação da utilização de serviços saúde ou determinação do número de visitas médicas, uma vez que se constituem em modelos para dados de contagem.

Também para estimação de dados de contagem Mullahy (1986) sugeriu a estimação do modelo “Hurdle Poisson” e Lambert e Greene (1992) propuseram o chamado “zero-inflated Poisson regression” (*ZIP models*). Zorn, em 1996, avaliou que os modelos Hurdle Poisson e o zero-inflated Poisson produzem estimativas bem parecidas, citando, inclusive, que a escolha entre um ou outro deve ser baseada em condições circunstanciais e de operacionalização.

Como já ressaltado, a estimação dos modelos com barreira pressupõe estimar o modelo de utilização dos serviços de saúde em duas etapas. Na primeira etapa, estima-se a probabilidade dos indivíduos contatarem ou não o serviço de saúde e na segunda, a decisão de frequência considerando-se apenas a amostra de pessoas com utilização positiva. Nesta primeira etapa do processo de estimação especifica-se um modelo de probabilidade binária (logit, probit, etc.) para determinar se o indivíduo procurou ou não algum serviço de saúde. Por exemplo, pode-se estimar uma regressão logística para toda a amostra, que neste caso, inclui pessoas que acessaram os serviços de saúde e pessoas que não acessaram.

A estimação da segunda etapa pode ser realizada de diversas formas. Uma delas consiste em estimar a regressão para a decisão de frequência pelo método dos mínimos quadrados ordinários (MQO). A fragilidade desse método é que os

dados de utilização dos serviços de saúde são censurados. Ao estimar por MQO essa particularidade da amostra é ignorada.

Um outro procedimento para esta segunda etapa consiste se estimar um modelo binomial negativo truncado ao zero para parte da amostra de interesse (por exemplo, para aqueles que realizaram consulta) que possibilite uma apuração mais adequada dos resultados.

Gerdtham (1997) e Pohlmeier e Ulrich (1994) empregaram essa metodologia para testar a hipótese de equidade horizontal no acesso aos serviços de saúde da Suécia e Alemanha respectivamente, considerando apenas a população adulta. Nos dois estudos, a decisão de contato e a decisão de frequência são determinadas por processos estocásticos distintos, sugerindo que a estimação do modelo *hurdle* é mais apropriada. Outros autores, como Viegas (2000) também utilizaram esta metodologia em seus trabalhos. Cameron *et al* (1988), é mais um exemplo de autor que estimou uma equação de utilização dos serviços de saúde para Austrália a partir do modelo binomial negativo para verificar a frequência com que os indivíduos utilizavam os serviços de saúde.

5.5.

O Modelo hurdle binomial negativo

Quando da abordagem do indicador de risco moral – IRM, foram verificados indícios da presença deste fenômeno no mercado de planos de saúde brasileiro. O IRM nos levou a pressupor a existência do risco moral quando foram analisados dados no nível nacional de abrangência.

Entretanto, as informações obtidas a partir dos resultados do indicador IRM podem ter alguns problemas, que serão tanto maiores quanto for o grau de afastamento da realidade das suposições consideradas na análise. Por exemplo, se pessoas que têm planos de saúde são pessoas que mesmo antes de ter a cobertura do plano já eram preocupadas com a saúde mais do que as que não têm plano, então pode ser que uma parcela do número de consultas médicas realizadas por essas pessoas decorra deste fato, e não do risco moral. Pode ser também que pessoas mais pobres talvez se consultem menos que pessoas com plano em função de associação entre ter plano e ter alto rendimento mensal, associação essa que pode fazer com que as restrições orçamentárias dos indivíduos mais pobres os

impecem de realizar consultas médicas de forma ideal, gerando um viés de alta no resultado do IRM.

Nesta parte da dissertação, o objetivo é construir um modelo econométrico que permita estimar a diferença entre o número de consultas de indivíduos com plano e o número de consultas de indivíduos sem plano controlando alguns possíveis fatores geradores de vícios nas estimativas. Para isso, é preciso que a análise do risco moral pela comparação do número de consultas ambulatoriais de pessoas “com plano” e pessoas “sem plano” leve em consideração algumas particularidades importantes que ocorrem em economia de saúde.

Uma característica importante a ser considerada por um modelo econométrico aplicado a tal situação é que a distribuição do número de consultas assume somente valores inteiros e não negativos. Desse modo, o modelo Poisson poderia se apresentar como sendo um ponto de partida interessante para estimação do número de consultas.

Sob a hipótese da distribuição Poisson, a probabilidade da variável aleatória NC (número de consultas) ocorrer n vezes, $P(NC=n)$ pode ser expressa pela seguinte equação:

$$P(y_i = n) = (e^{-\lambda} \lambda^n) / n! \quad \text{para } n = 0, 1, 2, \dots$$

onde,

$y_i = NC_i$ é o número de consultas do indivíduo i

λ é a média da distribuição, que pode ser expresso em forma de função dos atributos observados dos indivíduos.

Para incorporação de variáveis exógenas ao modelo, provenientes de um vetor p -dimensional $\mathbf{Z}$ de características do indivíduo i , basta substituir λ por alguma função de variáveis explicativas. Por exemplo, se o número de consultas é função de variáveis $(z_1, z_2, z_3, \dots, \in \mathbf{Z})$, tais como renda, sexo, idade, etc, então, pode-se fazer:

$$\lambda_i = \exp(\sum \beta_j z_{ij}) = \exp(\mathbf{z}'_i \mathbf{b})$$

onde,

β_j representa o coeficientes da j -ésima variável explicativa e

$\mathbf{z}$ é o vetor de variáveis explicativas do modelo.

Porém, existem características que não são incorporadas ao modelo que contém $Z'\beta$. São as variáveis desconhecidas e que contribuem para a soma dos quadrados dos erros da regressão. Essas variáveis, não incluídas no modelo por serem de difícil mensuração, formam o que se chama de heterogeneidade não observada. Com a finalidade de se controlar essa heterogeneidade não observada é necessário considerar λ como tendo um termo aleatório. Caso o efeito da heterogeneidade não observada seja descartado, o modelo pode apresentar dados sobre dispersos, ou seja, variância maior que a média. A sobre dispersão, por sua vez, gera problemas para estimação via Poisson, já que este modelo supõe que o valor esperado é igual à variância. A estimação do parâmetro λ em tais condições leva a estimativas ineficientes (Cameron e Trivedi (1986)).

Para contornar o problema, deve-se incluir termo estocástico (v_i) no modelo anterior, de forma que o mesmo passa a ser:

$$\lambda_i = \exp(\Sigma \beta_j Z_{ij}) \cdot \exp(v_i)$$

onde,

β_j representa o coeficiente da j -ésima variável explicativa,

Z é o vetor de variáveis explicativas do modelo e

v_i é o termo de erro aleatório.

O modelo binomial negativo é obtido ao se supor que este termo estocástico tem distribuição gama. Assim, com a incorporação do componente aleatório e^{v_i} , onde v_i tem distribuição gama, chega-se à especificação do modelo binomial negativo, dada por:

$$h(NC | l, a) = \frac{\Gamma(a^{-1} + NC)}{\Gamma(a^{-1})\Gamma(NC + 1)} \left(\frac{a^{-1}}{a^{-1} + l} \right)^{a^{-1}} \left(\frac{l}{l + a^{-1}} \right)^{NC}$$

A média e variância deste modelo são:

$$E[NC | l, a] = l$$

$$V[NC | l, a] = l(1 + a l)$$

onde NC é o número de consultas e α representa o termo de sobre dispersão dos dados. Quando $\alpha = 0$, a variância se iguala à média, indicando que não há sobre dispersão. Neste caso, o modelo Poisson poderia ser usado.

Olhando-se os dados da variável NC na tabela 17, observa-se que sua variância, à princípio, é estatisticamente maior que a média. O número de consultas apresenta média de 1,48 e variância de 8,11. O teste desta sobredispersão será apresentado mais adiante.

Tabela 18 – Distribuição do número de consultas

NC	Distribuição de NC	NC	Distribuição de NC
0	53,50	15	0,18
1	16,44	16	0,02
2	11,74	17	0,01
3	6,63	18	0,02
4	3,58	19	0,00
5	2,32	20	0,15
6	1,74	21	0,00
7	0,45	22	0,01
8	0,75	23	0,00
9	0,27	24	0,04
10	1,10	25	0,02
11	0,08	26	0,00
12	0,81	27	0,00
13	0,02	28	0,00
14	0,03	30	0,04
Media	1,483		
DP	2,84		
Variância	8,111		

Fonte: PNAD/1998.

O parâmetro α será testado e, caso seja verificada a sobre dispersão, a aplicação do modelo Poisson, como já enfatizado, fará com que os coeficientes sejam consistentes porém, os desvios padrões deixarão de ser.

Entretanto, existe ainda um outro problema para utilização do modelo Poisson: é a dependência nos dados. Para o caso aqui apresentado isto que dizer que a probabilidade de um indivíduo realizar uma consulta médica é dependente

do número de visitas que ocorreram anteriormente. O modelo Poisson supõe que as observações sejam independentes e identicamente distribuídas (iid).

Por todas esses comentários, o modelo binomial negativo é o mais adequado para ser aplicado a situação que se apresenta, onde se deseja estimar o número de consultas dos indivíduos em função de covariáveis. A estimação do modelo binomial negativo normalmente é feita por máxima verossimilhança e fornece estimativas consistentes e eficientes.

Deve-se levar em consideração para avaliação que o número de consultas é tal que o mesmo ocorre em função de dois processos distintos (Manning et al. 1981). O primeiro desses processos está relacionado a tomada de decisão para realização da primeira consulta sobre algum problema de saúde. Essa decisão é tomada invariavelmente pelo paciente. Já o número de consultas que se sucedem à primeira consulta – frequência – é em grande parte de responsabilidade dos médicos que o paciente consultou.

No caso do risco moral, espera-se poder avaliar sua ocorrência em função da estimação do número de consultas por intermédio de uma modelagem que englobe duas etapas distintas: uma equação relacionada a probabilidade de que um indivíduo realize consulta médica, estimada para o conjunto O ($O = O_1 + O_2$) da população, onde,

O representa o conjunto de indivíduos estudados, incluindo os que realizaram consulta médica e os que não realizaram;

O_1 representa o conjunto de indivíduos que realizaram consulta médica.

O_2 representa o conjunto de indivíduos que não realizaram consulta médica

E uma outra equação, relacionada a probabilidade do número de consultas dos indivíduos dado que, ao menos, uma consulta foi realizada. Aplicando-se este procedimento a dois grupos distintos de indivíduos, quais sejam, pessoas “com plano” e pessoas “sem plano”, pode-se chegar ao resultado desejado e conhecer a magnitude do risco moral nos dois processos.

Para estimação da primeira equação, pode-se usar uma regressão logística, com variável resposta binária, para avaliação das variáveis que interferem na probabilidade de uma pessoa realizar uma consulta médica. Esta equação tem a seguinte forma:

$$\Pr(C > 0, O) = 1 / (1 + \exp(-Z' \beta)),$$

onde,

- $\Pr(C, O)$ é a probabilidade do indivíduo i realizar consulta médica considerando-se o estudo sobre o conjunto O da população.
- β 's são os coeficientes a serem estimados para o i -ésimo regressor
- Z 's formam o conjunto de covariadas do modelo associado ao indivíduo i

Na segunda etapa do procedimento estima-se um novo modelo para determinação do número médio de consultas somente para indivíduos que realizaram consultas médicas. Devido a esta restrição, a modelagem é denominada de hurdle binomial negativo (ou, binomial negativo com barreira).

Após a estimação do modelo hurdle binomial negativo, pode-se analisar, a partir dos coeficientes das variáveis explicativas, a importância de cada uma nas probabilidades tanto de tomada de decisão sobre realizar consulta médica (primeira parte) quanto de frequência de consultas (segunda parte). Neste caso, a variável referente ao fato da pessoa ter plano de saúde seria uma das variáveis explicativas do modelo e seu coeficiente pode informar sobre a mudança na probabilidade de consulta ou frequência em função de se passar da situação de ter cobertura de plano de saúde para a de não ter.

5.5.1. Passos metodológicos para estimação do modelo

Os passos metodológicos para estimação do modelo hurdle binomial negativo passam pela definição de uma variável dicotômica d_i , que assume valores “1” ou “0” em função do indivíduo realizar ou não consulta médica.

$d_i = 1$, realiza consulta
0, caso contrário

de forma que:

$$\Pr(d_i = 1) = f_{1i} \quad \text{e} \quad \Pr(d_i = 0) = 1 - f_{1i} \quad i = 1, 2, \dots, n$$

onde $\Pr(d_i = 1) = f_{1i}$ é a probabilidade de d_i ser igual a 1

E f_{1i} pode ser modelado por regressão logística, tal como:

$$f_{1i} = \frac{1}{(1 + e^{-bZ})}$$

Se a y representar a variable aleatoria NC (número de consultas), então:

$$\begin{aligned} P(y_i) &= \Pr[d_i = 1, (y / y > 0)]^{d_i} * \Pr[d_i = 0, y = 0]^{1-d_i} & d_i = 0, 1 \\ &= \Pr[d_i = 1, (y / y > 0)]^{d_i} * \Pr[d_i = 0]^{1-d_i} \end{aligned}$$

onde $y = NC, \quad y = 0, 1, 2, 3, \dots$

Mas,

$$\begin{aligned} \Pr[d_i = 1, (y / y > 0)] &= \Pr[d_i = 1] * \Pr(y / y > 0) \\ &= f_{1i} * \Pr(y / y > 0) \end{aligned}$$

e $\Pr(y / y > 0)$ é uma função de probabilidade truncada em zero para dados de contagem. Se $f_2(y)$ for a distribuição para dados de contagem, como por exemplo a distribuição binomial negativa, tem-se que:

$$\Pr(y / y > 0) = \frac{f_2(y)}{(1 - f_2(0))}$$

$$\sum_{y=1}^{\infty} \frac{f_2(y)}{(1 - f_2(0))} = \frac{1}{1 - f_2(0)} \sum_{y=1}^{\infty} f_2(y) = \frac{1 - f_2(0)}{1 - f_2(0)} = 1$$

Este resultado decorre do fato de que

$$\sum_{y=0}^{\infty} f_2(y) = 1 = f_2(0) + \sum_{y=1}^{\infty} f_2(y) \quad \text{e} \quad \sum_{y=1}^{\infty} f_2(y) = 1 - f_2(0)$$

Assim,

$$\Pr(y) = \left[f_{1i} * \frac{f_2(0)}{1 - f_2(0)} \right]^{d_1} [1 - f_{1i}]^{1-d_1}$$

a partir de onde a função de verossimilhança pode ser obtida,

$$L(y) = \prod_{i=1}^n P(y_i)$$

$$L(y) = \prod_{i=1}^n \left[f_{1i} \frac{f_2(y)}{1 - f_2(0)} \right]^{d_i} [1 - f_{1i}]^{1-d_i}$$

e tomando-se o log, tem-se:

$$\log L(y) = \sum_{i=1}^n d_i \left[\log f_{1i} + \log \frac{f_2(y)}{1 - f_2(0)} \right] + \sum_{i=1}^n (1 - d_i) \log(1 - f_{1i}), \text{ com}$$

$$f_{1i} = \frac{1}{1 + e^{-bZ}}$$

se $f_2(y)$ tem distribuição binomial negativa, então:

$$\frac{f_2(y)}{1 - f_2(0)} = \frac{NB}{1 - NB(y=0)}, \text{ onde NB representa expressão da distribuição}$$

binomial negativa com parâmetros l e a , resultando em:

$$= \frac{\frac{\Gamma(a^{-1} + NC)}{\Gamma(a^{-1})\Gamma(NC + 1)} \left(\frac{a^{-1}}{a^{-1} + l} \right)^{a^{-1}} \left(\frac{l}{l + a^{-1}} \right)^{NC}}{1 - \left(\frac{\Gamma(a^{-1} + NC)}{\Gamma(a^{-1})\Gamma(NC + 1)} \left(\frac{a^{-1}}{a^{-1} + l} \right)^{a^{-1}} \left(\frac{l}{l + a^{-1}} \right)^{NC}, y = 0 \right)} \quad \text{para}$$

$$NC = 0, 1, 2, \dots$$

A função de verossimilhança para o processo é:

$$\log L(\Psi) = \sum_{i=1}^n (d_i \cdot \log f_{1i} + (1 - d_i) \log(1 - f_{1i})) +$$

$$\sum_{i=1}^n d_i [\log f_2(y) - \log[1 - f_2(0)]],$$

onde

$$\sum_{i=1}^n (d_i \cdot \log f_{1i} + (1 - d_i) \log(1 - f_{1i})) = \log L_1(\Psi) \text{ e}$$

$$\sum_{d_i=1} [\log f_2(y) - \log[1 - f_2(0)]] = \log L_2(\Psi)$$

onde,

- $\log L_1(\Psi)$ é o log de verossimilhança para um modelo logístico estimado para a amostra completa (incluindo pessoas que realizaram consulta e pessoas que não realizaram consulta).
- $\log L_2(\Psi)$ é o log de verossimilhança para um modelo binomial negativo truncado, estimado apenas para as pessoas que realizaram consulta.

A aplicação dos procedimentos para estimação do modelo hurdle binomial negativo aos dados da PNAD foi feita com o auxílio do software STATA.

5.5.2. Testes de especificação do modelo

Um primeiro procedimento é verificar a sobredispersão dos dados. Para isso, efetua-se o teste da razão de verossimilhança, comparando-se as diferenças entre os logaritmos da verossimilhança dos modelos Poisson e binomial negativo:

$$LR = -2 (LN_{Poisson} - LN_{negbin})$$

Onde:

$LN_{Poisson}$ = logaritmo da verossimilhança do modelo Poisson e

LN_{negbin} = logaritmo da verossimilhança do modelo binomial negativo

que permite-nos testar a seguinte hipótese:

$$H_0 : a = 0$$

$$H_1 : a > 0$$

Se a hipótese nula não for rejeitada, o modelo binomial negativo se reduz ao de Poisson, o que indica que os dados não contêm sobre dispersão. A hipótese nula será rejeitada quando a razão de verossimilhança for significativamente diferente de zero aos níveis usuais de análise estatística. O pacote STATA fornece o teste da razão de verossimilhança.

Para o modelo binomial negativo o log da verossimilhança foi -259300000. Para o modelo Poisson o valor correspondente foi -362000000. O teste da razão de verossimilhança apresentou p-valor abaixo de 0,0001, o que mostra que há sobredispersão. Os resultados estão no anexo 4.

O segundo estágio de análise engloba o teste do modelo binomial negativo contra o hurdle binomial negativo, o que nos permite verificar se o procedimento de “barreira”, com o objetivo de separar os dois processos decisórios (decisão de consulta e frequência) é eficiente.

Novamente, o teste da razão de verossimilhança é utilizado, agora para testar as hipóteses:

$$H_0 : b_1 = b_2$$

$$H_1 : b_1 \neq b_2$$

onde

b_1 é o vetor de parâmetros da etapa 1 (decisão de se consultar) e

b_2 é o vetor de parâmetros da segunda etapa (frequência).

Quando a hipótese nula é rejeitada, o modelo binomial negativo truncado (hurdle binomial negativo) é adequado e deve ser aplicado. Logo, o teste da razão de verossimilhança deve englobar a regressão logística e o modelo negativo binomial truncado para comparações com o negativo binomial. O teste tem a seguinte formulação:

$$LR_{\text{etapa final}} = 2(LN_{\text{negbin}} - LN_{\text{logit}} - LN_{\text{hurdle negbin}})$$

Este teste também é fornecido pelo STATA. Os resultados são descritos na sequência.

5.5.3. Interpretação dos Resultados

O teste da razão de verossimilhança para toda a amostra, feito com auxílio do STATA, comparando o modelo negativo binomial com o Poisson, a fim de testar a sobredispersão dos dados, indicou significância estatística para o parâmetro α , como pode ser visto na tabela 19.

Tabela 19 - Estimativa do parâmetro de sobredispersão e limites inferior e superior para significância de 0,05

Parâmetro	Estimativa	desv. Padrão	Lim. Inf.	Lim. Sup.
α	1.574795	.00031	1.574188	1.575403

Como indica a tabela 19, o intervalo de confiança para α indica que o mesmo é significativo para um nível de significância de 0,05. Isto significa que o modelo Poisson não é adequado para o caso que aqui se apresenta, pois há sobredispersão. A variância condicional do número de consultas é maior que a expectância condicional. Assim, as suposições consideradas para expectância e variância são:

$$E[NC | l, \alpha] = l \quad e$$

$$V[NC | l, a] = l(1 + a l)$$

Vale destacar que o modelo Poisson, como já enfatizado, também exigiria a suposição de que as consultas realizadas fossem processos independentes e identicamente distribuídos (iid), o que não é verdade. Logo, optou-se pela estimação via modelo hurdle binomial negativo (em duas etapas).

Com as saídas do STATA para a regressão logística (primeira etapa) é possível perceber que a decisão de realizar a primeira consulta é influenciada pela variável “PLANO”. Todas as variáveis foram significativas ao nível de significância de 0,01 e os odds ratio calculados são os tabelados abaixo:

Tabela 20 - Razão de chances (odds ratio) para regressão logística (primeira etapa)

Efeito	Odds Ratio	Significativo a 95% ?
_Isexo_2	1.809753	Sim
_Ieduca_2	.7036567	Sim
_Ieduca_3	.8678043	Sim
_Ieduca_4	.9842931	Sim
_Iidade_2	.7265672	Sim
_Iidade_3	.7268209	Sim
_Isaude_2	2.468592	Sim
_Isaude_3	3.961012	Sim
_Idcronica_2	.432314	Sim
_Irenda_2	1.194394	Sim
_Irenda_3	1.313569	Sim
_Irenda_4	1.493676	Sim
_Iplano_2	.4029656	Sim

Fonte: PNAD/1998

Obs: O nível de referência é o omitido

Pelos odds ratio da primeira etapa para variável “PLANO” nota-se que a mesma é significativa. Este resultado indica que a decisão de realizar a primeira consulta, tomada pelo paciente, muda de intensidade em função do indivíduo ter ou não cobertura por plano de saúde. As pessoas com plano de saúde têm probabilidades 60% maiores de consultar o médico do que pessoas sem cobertura de plano.

As mulheres consultam-se mais que os homens. Os odds ratio estimados indicam que as mulheres têm probabilidade maiores de ir ao consultório médico do que os indivíduos do sexo masculino.

Com relação à idade, constata-se que adultos e idosos consultam-se mais do que crianças. Entretanto, não foram observadas diferenças significativas entre o grupos de idades acima de 14 anos (grupo 2, 15 a 50 anos e grupo 3, acima de 51 anos).

A análise da variável renda indica que as pessoas com níveis mais elevados de renda tendem a realizar mais consultas do que pessoas com rendimentos menores. Os indivíduos com renda familiar per capita acima de R\$ 500,00 tomam a decisão de realizar consultas com probabilidade cerca de 40% maiores que pessoas com renda familiar per capita de até 100 reais.

Os resultados da segunda etapa, relacionados à frequência de realização de consultas, encontram-se na tabela abaixo.¹¹:

Tabela 21 - Estatísticas dos efeitos para segunda etapa: modelo negativo binomial com barreira

Efeito	Coef.	Std. Err.	z	P_valor	Interv. conf. 95%	
_Isexo_2	.29384	.0003184	922.85	0.000	.2932159	.2944641
_Ieduca_2	-.2087437	.0004293	-486.21	0.000	-.2095851	-.2079022
_Ieduca_3	-.1917292	.0006164	-311.04	0.000	-.1929374	-.1905211
_Ieduca_4	-.2034738	.0009228	-220.50	0.000	-.2052824	-.2016653
_Iidade_2	-.0775452	.0004815	-161.05	0.000	-.0784889	-.0766015
_Iidade_3	-.0830044	.0005722	-145.07	0.000	-.0841258	-.0818829
_Isaude_2	.5336285	.0004146	1287.17	0.000	.532816	.5344411
_Isaude_3	1.067756	.0007078	1508.56	0.000	1.066369	1.069143
_Idcronica_2	-.4016652	.0003775	-1064.14	0.000	-.402405	-.4009254
_Irenda_2	.0632484	.0004036	156.70	0.000	.0624572	.0640395
_Irenda_3	.0993329	.0004406	225.47	0.000	.0984695	.1001964
_Irenda_4	.1252742	.0005873	213.31	0.000	.1241231	.1264252
_Iplano_2	-.3956006	.0003875	-1020.94	0.000	-.3963601	-.3948412
_cons	1.061171	.0006413	1654.81	0.000	1.059914	1.062428

¹¹ Cabe ressaltar que nesta etapa não foi possível incorporar os efeitos do plano amostral devido à limitações do STATA

Estes resultados dos coeficientes, transformados em odds ratios, mostram que a variável “PLANO” também é importante na frequência das consultas médicas. O odd ratio estimado para os indivíduos sem plano é de 0,673 (igual a $\exp(-0,3956006) = 0,673$). Isto quer dizer que as pessoas com plano de saúde, têm um número esperado de consultas medicas cerca de 33% maior do que o número médio esperado de consultas de pessoas sem plano quando todas as demais variáveis são mantidas constantes.