



Universidade Estadual Paulista
Faculdade de Ciências e Letras
Departamento de Economia

MONOGRAFIA
Curso de Ciências Econômicas

**TEORIA DO DESENHO DE MECANISMOS: APLICAÇÕES NA
TEORIA DOS CONFLITOS INTERNACIONAIS**

Graduando: Guilherme de Abreu Sigrist

Orientadora: Prof. Dr^a. Claudia Heller

Banca Examinadora: Prof. Dr. André Luiz Correa

Araraquara

2011

RESUMO

O presente trabalho busca apresentar os conceitos principais referentes à Teoria do Desenho de Mecanismos e sua aplicação à Teoria dos Conflitos Internacionais, e como sua utilização pode solucionar alguns problemas referentes à metodologia mais utilizada na análise dos conflitos, baseada na Teoria dos Jogos tradicional.

Através de uma análise *game-free*, estabelecem condições gerais para a resolução de conflitos entre países, sem a necessidade de estabelecer formas de jogo específicas. Desse modo, determinam-se resultados para qualquer equilíbrio de qualquer forma de jogo que envolva países em disputa.

O trabalho segue apresentando a relação causal entre incerteza e incentivos para revelar de maneira desonesta informações privadas como causas fundamentais de guerra. Analisam-se separadamente as condições que levam à guerra em dois tipos de incerteza, constantemente discutidas na literatura de conflitos internacionais: incerteza em relação ao custo de guerra e em relação à força relativa. Os resultados são apresentados com a devida formulação matemática.

Sumário

INTRODUÇÃO.....	4
CAPÍTULO 1: A TEORIA DO DESENHO DE MECANISMOS	6
CAPÍTULO 2: A TEORIA DOS CONFLITOS INTERNACIONAIS	15
2.1) Conceitos Introdutórios	15
2.2) <i>Game-free analysis</i> e a Teoria do Desenho de Mecanismos	19
CAPÍTULO 3: MODELO DE <i>GAME-FREE</i> PARA CONFLITOS INTERNACIONAIS	22
3.1) Resultados e Formulações Iniciais.....	22
3.2) Análise das Incertezas	30
(3.2.1) Incerteza em Relação aos Custos de Guerra	31
(3.2.2) Incerteza em Relação à Força Relativa	38
CONCLUSÕES.....	45
BIBLIOGRAFIA	47

1) INTRODUÇÃO

Desde meados da segunda metade do século XX, a teoria dos mecanismos, ramo da tradicional teoria dos jogos, vem sendo base teórica para a análise de diversas questões dentro da economia, principalmente as que dizem respeito às instituições de mercado e como alocam os recursos, dado o diferencial de informações entre os indivíduos.

A importância da teoria dos mecanismos dentro da comunidade científica pode ser comprovada pelo Nobel concedido a Leonid Hurwicz, Eric Maskin, e Roger Myerson em 2007, por terem fundamentado as bases da teoria.

A principal contribuição da teoria é a de permitir a elucidação de problemas cuja característica é a diferença de informações entre os agentes e a necessidade de incentivá-los a revelar de maneira adequada suas informações. Através de seus conceitos, que serão desenvolvidos ao longo do trabalho, é possível que, através do mecanismo adequado, os agentes, mesmo agindo de acordo com seu próprio interesse, revelem suas informações honestamente aos outros, acreditando que os outros farão o mesmo. Dessa forma, a possibilidade de estabelecer acordos, tomando o pressuposto que os indivíduos agirão de maneira honesta, se torna muito mais provável. A teoria dos mecanismos torna-se, nesse contexto, uma forma de estabelecer equilíbrios onde a diferença de informação é uma das limitações para a definição de acordos.

A dificuldade de se formar acordos por falta de informação pode ser observada, de maneira específica, em relação a países que se encontram em plena eminência de um conflito. A dificuldade de monitorar as informações uns dos outros dificulta o estabelecimento de acordos. A teoria dos conflitos internacionais tem como principal fundamento analisar esse tipo de questão, em especial os problemas que definem um conflito ou um acordo de paz. Questões centrais na teoria dos conflitos internacionais enfocam, como em Mitchell (1994), a dificuldade em construir instituições capazes de aliar os objetivos dos países em disputa e reduzir conflitos desnecessários.

A utilização da teoria dos mecanismos na teoria dos conflitos internacionais pode facilitar a determinação de acordos pacíficos, dado que as partes (indivíduos ou países) revelarão informações de maneira honesta.

O objetivo do trabalho é descrever como os autores conciliam as ideias das duas teorias na busca de soluções que descrevam as possibilidades de embate direto entre as nações ou a formação de acordos de paz.

CAPÍTULO 1: A TEORIA DO DESENHO DE MECANISMOS

De acordo com Myerson (1988), um mecanismo é uma especificação de como as decisões dentro da economia são determinadas como funções das informações detidas pelos indivíduos. Analisando desta maneira, permite-se que uma vasta gama de instituições e organizações econômicas possa ser analisada dentro de um mecanismo, dotado de informações. Desse modo, a teoria dos mecanismos é capaz de oferecer uma estrutura em que diversas instituições podem ser comparadas e encontrar as mais eficientes.

O princípio básico da teoria dos mecanismos é o de que os incentivos têm restrições que devem ser analisadas da mesma maneira que as dos recursos (escassos) dentro da economia. Para melhor exemplificar, Myerson (1988) afirma que em situações em que é difícil monitorar as ações dos indivíduos, a necessidade de oferecer incentivos aos agentes para compartilharem suas informações ou colaborarem de alguma outra forma impõe as mesmas restrições que a escassez de recursos (matérias-primas, por exemplo).

A preocupação com a definição desses tipos de mecanismos começou, segundo Myerson (1988), com o estudo de outro tipo de mecanismo chamado *direct-revelation mechanism*. Esse mecanismo toma por base um mediador que pode, de maneira confidencial, comunicar-se separadamente com cada indivíduo da economia. O mediador pode ser considerado uma pessoa confiável e com bom conhecimento sobre os fatores econômicos envolvidos nas decisões dos agentes ou um computador capaz de computar as diferentes informações adquiridas junto aos indivíduos. Definido este princípio, a cada estágio do processo econômico, os agentes reportam ao mediador todas as suas informações privadas, que compreendem tudo o que eles sabem que os outros agentes envolvidos no processo podem não saber. Então, a partir das informações recebidas individualmente dos indivíduos, o mediador recomenda (confidencialmente também) a cada indivíduo uma ação. O *direct-revelation mechanism* é, portanto, qualquer regra ou conjunto de regras que definem as recomendações do mediador como funções das informações reveladas pelos agentes.

Outro conceito relevante para a teoria dos mecanismos é chamado *incentive compatibility*, originalmente elaborado por Leonid Hurwicz (1972). Segundo o autor, o conceito de *incentive compatibility* permite descrever um conjunto de procedimentos (ou mecanismos) pelos quais os indivíduos, baseados em seus próprios interesses, agem de maneira não estratégica, ou em outras palavras, de maneira verdadeira, confiável. Assim, um mecanismo é dito *incentive-compatible* quando todos os indivíduos dentro do processo econômico acreditam que os outros reportarão de maneira honesta ao mediador e serão obedientes às suas recomendações. Sendo assim, nenhum indivíduo esperará melhores resultados agindo de maneira desonesta ou desobediente, dado o conjunto de informações que possui. Se agir de maneira honesta e obediente configura um equilíbrio no sentido definido pela teoria dos jogos, então o mecanismo que define essas ações é denominado *incentive-compatible*.

Deve-se reconhecer as limitações desse tipo de mecanismo, devido à inexistência de um mediador com uma ação tão ampla e centralizada como o apresentado e também à observação comum de casos de desonestidade dentro do sistema econômico.

Entretanto, de acordo com Myerson (1988), a importância do estudo desse tipo de mecanismo (agora denominado *incentive-compatible direct-revelation mechanism*) é derivada da ideia de que para qualquer equilíbrio de qualquer mecanismo, há um *incentive-compatible direct-revelation mechanism* equivalente. Desse modo, dada a equivalência com qualquer equilíbrio de qualquer mecanismo, é possível aferir o que pode ser alcançado em todos os possíveis equilíbrios de todos os possíveis mecanismos, para uma situação econômica dada. Em outras palavras, qualquer resultado de equilíbrio de um mecanismo tomado arbitrariamente pode ser reproduzido por um *incentive-compatible direct-revelation mechanism* (Nobel 2007). Esse conceito é denominado *revelation principle*, e foi sintetizado por Myerson (1982) nas situações onde a honestidade e a obediência são limitações impostas à teoria.

Os conceitos por trás do *revelation principle*, define Myerson (1979), são os seguintes. Primeiramente, um mediador que coletou todas as informações relevantes detidas pelos indivíduos na economia pode simular o retorno de cada indivíduo ou instituição e então lhes oferecer recomendações. As ações recomendadas pelo mediador são ações que os

agentes poderiam ter em qualquer outro tipo de mecanismo (daí a equivalência com outros mecanismos). Segundo, quanto mais informação um indivíduo tem, mais difícil poderá ser preveni-lo de descobrir algum meio de ganhar desobedecendo ou enganando o mediador. Então, a necessidade de oferecer incentivos aos indivíduos será menor, ou em outras palavras, a restrição acarretada pelos incentivos será enfraquecida quando o mediador revela para cada indivíduo apenas o mínimo de informação necessária para ele identificar a ação que lhe foi recomendada, e nada a mais sobre as recomendações dos outros indivíduos. Sendo assim, e definindo o mediador como um processador de informações confiável e verdadeiro, pode-se assumir que os agentes irão transmitir todas suas informações ao mediador, caracterizando o processo denominado *maximal revelation* ao mediador, enquanto o mediador transmitirá apenas o mínimo para que os indivíduos identifiquem sua recomendação, caracterizando *minimal revelation* aos indivíduos para o enfraquecimento das restrições ocasionadas pelos incentivos.

Se os indivíduos, dentro do *revelation principle*, revelarem de maneira honesta suas informações e as recomendações do mediador caracterizarem uma função das informações recebidas de cada um, se um indivíduo ganhar desobedecendo às recomendações do mediador, então ele ganhará desobedecendo a sua própria estratégia, pois ela, determinada pelas suas informações, é parte da função de recomendações do mediador, estabelecendo uma inconsistência lógica.

Na sequência deste primeiro capítulo será proposta a formulação matemática geral da prova do *revelation principle*, numa situação em que indivíduos detêm informações privadas as quais poderiam em eventual situação ser distorcidas ou negligenciadas, mas não há dúvidas quanto à obediência às ações ou práticas recomendadas pelo mediador. Myerson em diversos de seus textos (1979, 1982, 1988) demonstra este problema, os quais serão explorados a seguir.

Primeiramente, adota-se um número n de indivíduos, numerados a partir de 1 até n . Denota-se por C o conjunto de todas as possíveis combinações de ações ou alocações de recursos que os indivíduos dentro da economia possam escolher. De acordo com Harsanyi (1967), um dos precursores da teoria, os agentes detêm informações a respeito de suas próprias preferências e vontades, mas também reservam crenças a respeito das

informações que os outros indivíduos possuem e podem supor de determinada maneira o que eles venham a fazer. Num contexto geral, os indivíduos têm uma distribuição de probabilidades subjetiva em relação às diversas probabilidades dentro da economia. O autor propõe:

“(...)a new theory for the analysis of games with incomplete information where the players are uncertain about some important parameters of the game situation, such as the payoff functions, the strategies available to various players, the information other players have about the game, etc. However, each player has a subjective probability distribution over the alternative possibilities” (Harsanyi, 1967, p.159)

O autor define então a informação privada do indivíduo (em relação a ele mesmo e aos outros agentes) como *type* do indivíduo. O conjunto de todos os *types* para qualquer indivíduo i será denominado T_i , e a combinação de todos os *types* para todos os indivíduos será $T = T_i \times \dots \times T_n$.

Em relação às preferências¹ de cada indivíduo i , elas podem ser descritas como uma função de *payoff* que encontre dentre todas as possíveis combinações de ações ou alocação de recursos alguma combinação que seja baseada em suas crenças a respeito das ações ou das combinações de recursos que os outros indivíduos adotarão e em suas próprias convicções. Genericamente então, podem ser descritas como $u_i: C \times T \rightarrow \mathbb{R}$, sendo $u_i(c, (t_i, \dots, t_n))$ o *payoff* que o indivíduo i obteria caso as combinações de ações ou recursos c (alocação dentro do conjunto C) se realize e que os *types* de cada indivíduo 1 a n nesse processo sejam referentes a $(t_i, \dots, t_n)$. Fica então evidente que $t = (t_i, \dots, t_n)$ representa uma combinação de *types* de todos os indivíduos. (Myerson, 1982).

Em Myerson (1988), as crenças de cada indivíduo i , podem ser descritas de maneira genérica como uma função probabilidade $p(\cdot | \cdot)$, onde $p_i(t_i, \dots, t_{i-1}, t_{i+1}, \dots, t_n | t_i)$ é a probabilidade do indivíduo i acreditar que os *types* dos demais indivíduos ($n-i$), dado

¹ Cabe salientar que os *payoffs* são determinados em escala de utilidade de Neumann – Morgenstern, que como especifica Binmore (2009), baseia-se nos riscos que o indivíduo corre ao tomar determinada decisão e nas decisões dos outros indivíduos.

seu próprio *type* (t_i), serão iguais a $(t_i, \dots, t_{i-1}, t_{i+1}, \dots, t_n)$. Dessa maneira, $(t_i, \dots, t_{i-1}, t_{i+1}, \dots, t_n)$ representa uma combinação dos *types* de todos os indivíduos exceto i . De uma maneira geral, a expressão $T_{-i} = T_1 \times \dots \times T_{i-1} \times T_{i+1} \times \dots \times T_n$ descreve o conjunto de todas as combinações de *types* de todos os indivíduos exceto i .

O autor conclui que o modelo geral de uma economia definida por essas estruturas² é denominado *Bayesian collective-choice problem*.

Aumann (1967) já especificava a determinação de equilíbrios estruturados pela racionalidade bayesiana acima descrita. Cabe, primeiramente, identificar um mecanismo geral que permita a determinação de equilíbrio seguindo esta lógica.

Myerson (1979) analisa este tipo de mecanismo como uma função que permite combinar as estratégias de cada indivíduo i dentro do espaço econômico na determinação da escolha das ações e da alocação de recursos. Formalmente, dados os princípios do *Bayesian collective-choice problem*, um mecanismo geral seria qualquer função que se estabeleça na forma $\gamma: S_1 \times \dots \times S_n \rightarrow C$, onde, para cada i , S_i é um conjunto não vazio que denota o conjunto de estratégias possíveis dentro do espaço C para o indivíduo i .

Este mecanismo especifica como as alocações de recursos e as combinações de ações dentro da economia são definidas por qualquer combinação de estratégias escolhidas por cada indivíduo i dentro de uma gama de opções definidas por S .

Um equilíbrio, dentro deste tipo de mecanismo, é qualquer especificação de como cada indivíduo toma suas ações baseadas numa função de seu próprio *type*. Sendo essas ações tomadas desta maneira pelos n indivíduos dentro da economia, baseando apenas em suas próprias informações, isto é, pelos seus próprios *types* (cada indivíduo conhece suas preferências e infere sobre as dos outros, mas não tem conhecimento a respeito dos *types* dos outros indivíduos), nenhum indivíduo espera um resultado (*payoff*) melhor se

² Estrutura que comporte as variáveis C, T, u e p em todos seus níveis:

$(C, T_1, \dots, T_n, u_1, \dots, u_n, \dots, p_1, \dots, p_n)$

desviando unilateralmente do equilíbrio seja mentindo, escondendo informações, etc. (Harsanyi, 1967)

Sendo assim, em Myerson (1988), pode-se definir um equilíbrio como $\sigma = (\sigma_1, \dots, \sigma_n)$, referente ao mecanismo γ se para cada indivíduo i , σ_i é uma função que relacione o conjunto de *types* com o conjunto de estratégias possíveis deste indivíduo i , assim sendo, de T_i e S_i , para cada t_i em T_i e cada s_i em S_i .

O processo pode ser descrito como:

$$\sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\gamma(\sigma(t)), t) \geq \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\gamma(\sigma_{-i}(t_{-i}), s_i), t).$$

Sendo: $\sigma(t) = \sigma_1(t_1), \dots, \sigma_n(t_n)$

$$\sigma(\sigma_{-i}(t_{-i})) = \sigma_1(t_1), \dots, \sigma_{i-1}(t_{i-1}), s_i, \sigma_{i+1}(t_{i+1}), \dots, \sigma_n(t_n)$$

O que se pode observar nesta equação é que para cada *type* t_i dentro do conjunto T_i , seguindo a lógica do mecanismo anteriormente especificado, em que a alocação dos recursos (ou das ações) é determinada pelas estratégias de cada indivíduo, no caso em que um indivíduo i , que espera que todos os outros sigam a estratégia especificada pelo equilíbrio σ , altera sua estratégia de $\sigma_i(t_i)$ para s_i , o somatório dos *payoffs* de cada indivíduo será no máximo igual ao do equilíbrio proposto σ^3 . Em outras palavras, Dasgupta, Hammond e Maskin (1979) definem que, num equilíbrio σ , nenhum indivíduo i , conhecendo apenas seu próprio *type* t_i , pode aumentar seu *payoff* esperado

³ Pode-se observar sistematicamente que $u_i(\gamma(\sigma(t)), t) \geq u_i(\gamma(\sigma_{-i}(t_{-i}), s_i), t)$, o que evidencia que a mudança de estratégia de $\sigma_i(t_i)$ para s_i acarreta não mais que a identidade dos *payoffs*, nunca um aumento.

mudando de estratégia. Palfrey e Srivastava (1987) ainda chamam esse tipo de equilíbrio de “Equilíbrio Bayesiano”, pelo fato de respeitar a condição de que cada jogador, ao definir suas estratégias (em S_i), as toma apenas conhecendo seu próprio *type*.

Após a análise destas condições, pode-se voltar a alguns conceitos propostos no início do capítulo. De acordo com Myerson (1982), um *direct-revelation mechanism* (proposto anteriormente como qualquer regra ou conjunto de regras que definem as recomendações do mediador como funções das informações reveladas pelos agentes) pode ser definido como qualquer mecanismo em que o conjunto de estratégias possíveis do jogador i (S_i) é igual ao conjunto de todos os possíveis *types* para o mesmo jogador T_i .

Além disso, um *direct-revelation mechanism* é *incentive-compatible* se, e somente se, houver um equilíbrio bayesiano, como o formulado anteriormente, que faça sempre cada indivíduo reportar ao mediador suas informações verdadeiras, ou seja, seu *type* real.

Se o mecanismo anterior era proposto de maneira e relacionar o conjunto de estratégias de cada indivíduo i , sendo $\gamma: S_1 \times \dots \times S_n \rightarrow C$, a igualdade entre o conjunto de estratégias e de *types* definidas acima pelo *direct-revelation mechanism*, permite dizer que $\mu: T_1 \times \dots \times T_n \rightarrow C$ pode representar um *direct-revelation mechanism*.

Como visto na introdução deste capítulo, muitas vezes há a necessidade de oferecer incentivos aos agentes para que compartilhem suas informações ou colaborem de alguma outra forma. Isto leva às restrições denominadas restrições por incentivos. Indivíduos que necessitem de incentivos para revelar de maneira honesta suas informações ou problemas de seleção adversa, (assinalado por Kreps (1994), como quando uma ou mais das partes envolvidas numa transação detém informações relevantes, mas desconhecidas pelas outras partes, representando um caso de informação assimétrica) muitas vezes levam estes a tomarem mudarem suas estratégias, valendo-se de suas vantagens.

Num *incentive-compatible direct-revelation mechanism*, sempre é a melhor decisão (ou seja, os indivíduos não esperam *payoffs* maiores) para os indivíduos expressarem de maneira verídica suas informações. Em Myerson (1988):

$$\sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\mu(t), t) \geq \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\mu(t_{-i}, r_i), t).$$

Onde:

$$(t_{-i}, r_i) = (t_1, \dots, t_{i-1}, r_i, t_{i+1}, \dots, t_n)$$

O r_i refere-se ao *type* de um indivíduo que utiliza suas vantagens de informação, por algum dos motivos descritos acima. Neste sentido, o somatório dos *payoffs* dos indivíduos em um equilíbrio dentro de um *incentive-compatible direct-revelation mechanism*, quando há a necessidade de estabelecer incentivos ou quando indivíduos usam de maneira vantajosa suas informações, será no máximo igual ao somatório de quando os indivíduos agem de maneira honesta.

Dessa maneira, os indivíduos não esperam nada melhor não revelando de maneira adequada suas informações.

Cabe agora provar o *revelation principle*, que preconiza que qualquer resultado de equilíbrio de um mecanismo tomado arbitrariamente pode ser reproduzido por um *incentive-compatible direct-revelation mechanism* (Nobel, 2007). Desta maneira, dado qualquer mecanismo, já representado por γ , qualquer equilíbrio no sentido de Bayes σ dentro do mecanismo γ , um *direct-revelation mechanism* representado por μ para cada t em T , tem-se:

$$\mu(t) = \gamma(\sigma(t)).$$

Assim sendo, como afirma Myerson (1988) o mecanismo especificado por μ (*direct-revelation mechanism*) sempre levará à mesma alocação de recursos (ou em termos mais gerais, às mesmas escolhas sociais) que γ , quando os indivíduos agirem como no equilíbrio σ .

Além disso, μ é *incentive-compatible* porque para qualquer indivíduo i e qualquer types t_i, r_i em T_i :

$$\begin{aligned} \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\mu(t), t) &= \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\gamma(\sigma(t)), t) \\ &\geq \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\gamma(\sigma_{-i}(t_{-i}), \sigma_i(r_i)), t) \\ &= \sum_{t_{-i} \in T_{-i}} p_i(t_{-i} | t_i) u_i(\mu(t_{-i}, r_i), t). \end{aligned}$$

Sendo assim, μ é *incentive-compatible direct-revelation mechanism*, equivalente ao mecanismo γ e seu equilíbrio σ .

O que se pôde observar é que qualquer mecanismo e seus múltiplos equilíbrios são equivalentes a um *incentive-compatible direct-revelation mechanism*.

CAPÍTULO 2: A TEORIA DOS CONFLITOS INTERNACIONAIS

2.1) Alguns Conceitos

Uma das questões centrais a respeito dos conflitos internacionais é entender os obstáculos que impedem que países envolvidos em qualquer tipo de disputa possam alcançar acordos mutuamente benéficos.

Segundo Fey e Ramsay (2009), a pergunta que deve ser feita é por que países entram em longos e custosos conflitos sendo que poderiam resolver suas pendências através de acordos sem a necessidade de guerra, obtendo resultados mais satisfatórios.

A literatura sobre conflitos internacionais sempre trouxe como principal aspecto que impede arranjos benéficos em tempos de conflito a incerteza em relação aos adversários, originando incentivos que levam os países a revelarem de maneira desonesta suas informações, impedindo a formação de qualquer tipo de acordo (Fearon, 1995). A importância dessa literatura é a de identificar o diferencial de informação como inerente no processo de barganha e no comportamento de guerra de cada país. Devido às informações privadas, os países podem não ser capazes de alcançar acordos adequados, sendo então de suma importância conhecer mais a respeito de si mesmos. Entretanto, devido ao diferencial de informações e aos incentivos de não revelarem suas informações de maneira confiável, não há como os países acreditarem na veracidade de seus adversários, inviabilizando os acordos.

Em relação à análise teórica, na maior parte da literatura os modelos escolhidos para trabalhar com a questão da incerteza dentro da teoria dos conflitos internacionais foram aqueles preconizados pela teoria dos jogos convencional. De acordo com Fey e Ramsay (2009), esse tipo de abordagem gera alguns problemas quando aplicada a conflitos. O principal problema é que qualquer equilíbrio de um determinado jogo é tipicamente sensível às particularidades do jogo. Para se especificar um jogo, é necessário esclarecer todos os detalhes relevantes, como as escolhas possíveis para os jogadores, o tempo necessário para efetuar-las, as preferências e informações dos tomadores de decisões. Essa especificidade pode acarretar problemas em formulações mais gerais. O fato de, por exemplo, dentro de determinado jogo (representado por uma barganha ou acordo

entre países), as nações decidirem guerrear for explicado pela diferença de informações, nada prova que as outras possíveis variações deste mesmo jogo terão o mesmo resultado, ou a causa desse mesmo resultado (ir à guerra) seja explicada pelo diferencial de informações.

Em seu texto de 2004, Ramsay compara a formulação de modelos teóricos que tratam de processos eleitorais aos que tratam de conflitos internacionais. Segundo o autor, os problemas gerados na utilização da teoria dos jogos são maximizados quando a teoria é aplicada a conflitos. Nos processos eleitorais, as especificidades dos jogos costumam ser constantes, ou seja, geram jogos “normais”: candidatos apresentam sua candidatura, conduzem suas campanhas, os eleitores em determinado dia do ano os elegem de maneira simultânea, etc. Já na teoria dos conflitos, os processos são muito mais complexos. Países podem sugerir acordos ou entrar em guerra a qualquer momento, sem especificações prévias; podem começar uma guerra se sua proposta for recusada, ou ainda entrar em guerra mesmo quando aceita; o processo de barganha pode, ao invés de ser bi/multilateral, envolver algum tipo de mediador ou alguma outra entidade neutra. Dessa maneira, devido à necessidade de especificações, impede-se a aplicabilidade de alguma forma específica de jogo.

A abordagem elaborada por Fey e Ramsay (2009) enfatiza principalmente a dificuldade de se aplicar os conceitos da teoria dos jogos convencional na análise de eventos da teoria dos conflitos internacionais. A dificuldade em se obter um modelo específico que represente uma barganha em tempos de crise leva à necessidade de estabelecer uma aproximação que determine resultados para qualquer equilíbrio de qualquer jogo em que países tenham que estabelecer acordos frente à eminência de uma guerra.

Seria inviável analisar todos os possíveis jogos, devido o incomensurável número de possibilidades, e dentro das possíveis formas destes jogos, os diversos equilíbrios a serem obtidos. Para tanto, a Teoria do Desenho de Mecanismos aparece como uma saída para a solução desse problema. Dessa maneira, é possível analisar os resultados dos processos de barganha sem a necessidade de especificação das formas do jogo. Utilizando o *revelation principle* podem-se identificar equilíbrios onde acordos são determinados voluntariamente, evitando conflitos, sem a necessidade de estabelecer determinada forma de jogo.

Essa formulação, sem a necessidade de especificação de um modelo de barganha, foi descrita por Banks (1990) como “*game-free analysis*”, ou seja, sem a necessidade de escolha de um modelo de jogo.

Segundo o autor, esse tipo de análise permite identificar, de uma maneira geral, como o diferencial de informações pode limitar a formulação de acordos pacíficos e levar os países à guerra, mesmo quando há interesse comum em evitá-las.

Fearon (1995) ainda enfatiza que além da incerteza quanto aos adversários, os incentivos a revelarem de maneira não confiável suas informações causam grandes entraves nas barganhas em tempos de crise.

Como já analisado, através do *incentive-compatibility*, torna-se possível analisar os casos onde haverá equilíbrios em que os agentes não terão incentivos em revelar suas informações de maneira desonesta. De maneira análoga, permite identificar também casos onde o diferencial de informações e os incentivos para não compartilhar suas preferências se tornarão obstáculos intransponíveis para a realização de acordos pacíficos.

Fey e Ramsay (2009) foram além das constatações de Banks (1990) e Fearon (1995) quanto à incerteza como principal obstáculo no processo de barganha. Os autores identificaram que o sucesso ou não de uma barganha depende também do tipo de incerteza que os países enfrentam. No caso em que os países estão incertos quanto aos custos de seus adversários de entrarem em guerra, mas não tem incerteza quanto à probabilidade de sucesso no conflito, em qualquer equilíbrio de qualquer jogo, a probabilidade esperada de guerra e a utilidade esperada variam de maneira inversa (ou monotonicamente decrescente) em relação ao custo de guerra. Além disso, demonstra-se que nesse tipo de incerteza existem jogos em que a ocorrência ou não de guerra (ou o sucesso ou não de um acordo) pouco depende da diferença de informações entre as nações (expressados aqui pelo custo de guerra). Em outras palavras, as estruturas de *payoffs* das nações são praticamente insensíveis ao diferencial de informações.

Outro exemplo de incerteza é quando os valores dos países em entrarem em conflito são interdependentes. Este caso se dá principalmente quando há incerteza quanto à força relativa dos países, mas não há incerteza quanto aos custos. Neste exemplo, se os custos são baixos, incentivos para revelar de maneira não adequada suas informações e a

existência de diferencial de informação sempre gerarão probabilidades positivas de conflito. Cabe aqui então o papel de uma nação ou outra entidade internacional neutra de estabelecer incentivos que levem estes países a estabelecer acordos pacíficos. Estes processos mostram os dois lados da incerteza: enquanto a incerteza em relação aos custos pode ser “benigna” na determinação de uma barganha em tempos de crise, por não influenciar a decisão dos países diretamente, a incerteza quanto à força relativa gerará sempre mais problemas na obtenção de algum acordo.

A *game-free analysis* de Banks (1990), como modelo norteador da análise teórica dos conflitos internacionais ainda se diferencia da análise econômica tradicional do processo de barganha⁴ em outros diversos pontos. As disputas internacionais se dão, segundo Waltz (2001), por países soberanos que não estão submetidos a leis que regem as políticas entre eles, sendo suas decisões guiadas apenas por sua própria ambição. Posto desta maneira, um país pode escolher a qualquer momento ir à guerra, independente se o outro país queira o conflito ou não. Desse modo, o conflito se dará sem a necessidade de que ambos os lados estejam de comum acordo. Assim sendo, acordos só podem se dar de forma voluntária.

Esta constatação revela que, no âmbito internacional, em que os países não são obrigados a cumprir acordos realizados com os demais, o que difere da teoria convencional, que estabelece em acordos de barganha que, explícita ou implicitamente, haja o cumprimento dos contratos (Myerson e Satterhwaite, 1983) Em outras palavras, os agentes não podem voltar atrás em um contrato que foi previamente acordado.

Isso revela que, devido à falta de contratos exigíveis nas relações internacionais, os países podem impor veto em relação ao que foi definido a qualquer momento, tanto durante a elaboração quanto depois de estabelecido.

Segundo Schelling (1960), outra questão importante em relação à teoria dos conflitos internacionais é a de que os países podem aprender com seus oponentes durante o processo de negociação.

Isso implica que a decisão final de um país entrar ou não em guerra deve incorporar toda a informação adicional transmitida pelos seus oponentes durante o processo de negociação.

⁴ Aqui, os processos de conflito e de barganha serão bilaterais, mais comumente usados na literatura.

Fey e Ramsay (2009) ainda dão um passo seguinte em relação a *game-free analysis* de Banks (1990). A teoria originária de Banks considera a informação incompleta apenas de um lado do conflito. Desse modo, com essa estrutura de informação todos os equilíbrios de todos os jogos de barganha são monotônicos, ou seja, quanto mais informado é um país, maior será seu *payoff* ao estabelecer um acordo⁵.

2.2) *Game-free analysis* e a Teoria do Desenho de Mecanismos

Quando são utilizados modelos formais para analisar determinados eventos, os resultados obtidos a partir desta análise dependem fundamentalmente das suposições que foram utilizadas na formação destes modelos. De acordo com Mark e Fey (2008), o que é pouco discutido dentro da literatura é o quão sensível são estes resultados às peculiaridades que definem o modelo escolhido, ou mais especificamente, à forma de jogo escolhida.

Isso configura um grande entrave na hora de aplicar modelos: muitas vezes, os diversos resultados possíveis dentro dos diversos equilíbrios dos mais variados jogos podem ser muito sensíveis às diversas especificidades do modelo, tornando difícil a obtenção de uma generalidade ou de um padrão dentro de determinado fenômeno.

Dessa maneira, o que se pode perguntar é com que intensidade um jogo específico pode ser aplicável.

Fearon (1995) estabelece de maneira geral que:

“Under very broad conditions, if A cannot learn B's private information and if A's own costs are not too large, then state A's optimal grab produces a positive chance of war. Intuitively, if A is not too fearful of the costs of war relative to what might be gained in bargaining, it will run some risk of war in hopes of gaining on the ground.” (Fearon, 1995, pp. 194-195)

⁵ Dá-se pelo fato de que, o país com mais conhecimento pode definir os termos do contrato de maneira mais vantajosa.

Fearon estabelece nessa sentença uma generalidade em relação à teoria dos conflitos internacionais. Se o custo do país A for pequeno, então a diferença de informações entre os países A e B produzirá probabilidade positiva de guerra.

Entretanto, essa generalidade é fundamentada em um modelo em que a informação incompleta é em apenas um lado do conflito. Logo, essa generalidade pode não ser aplicável a qualquer tipo de processo de barganha, como afirmam Fey e Ramsay (2008, 2009).

Powell (2004) afirma que:

“(...) the potential sensitivity of informational accounts of war to the bargaining environment to the sources of uncertainty and the ability to resolve that uncertainty. This sensitivity means that future informational efforts should either include some robustness checks for a particular formalization of the bargaining environment or they should ground that formalization empirically. For example, a satisfactory explanation of prolonged fighting that only holds if the states cannot make offers quickly must also explain in an empirically convincing way why the states cannot make offers quickly.” (Powell, 2004, p. 352)

O autor revela a necessidade do embasamento teórico e empírico necessário para a validação da escolha de um modelo que explique determinado fenômeno. Entretanto, esses pressupostos adotados podem não ser verificados em todas as possíveis situações em que os países podem estar com seus adversários. Dessa maneira, modelos específicos pouco têm a acrescentar em formulações mais gerais.

Cabe então, na resolução deste dilema, a utilização da metodologia da Teoria do Desenho de Mecanismos, como já especificado anteriormente

Detalhando-se melhor a relação entre a Teoria dos Conflitos Internacionais e a Teoria do Desenho de Mecanismos, o que se pretende constatar é a possibilidade de análise dos *payoffs* sem a necessidade de procedimentos precisos e modelos limitados. Pode-se determinar quais são os possíveis resultados para todos os processos de barganha a serem utilizados.

Isso se dá pela capacidade da Teoria do Desenho de Mecanismos de prover resultados sem a necessidade de escolher algum jogo, possibilitando uma análise *game-free*. Como já dito, a utilização do *revelation principle* reduz a análise de todos os possíveis jogos para apenas os classificados dentro do *incentive compatible direct mechanism*.

Em específico à Teoria dos Conflitos Internacionais, isso significa que qualquer *payoff* resultante de qualquer equilíbrio de qualquer procedimento de barganha pode ser alcançado por um equilíbrio determinado por um *incentive compatible direct mechanism*, implicando que cada jogador revelará suas informações de maneira confiável e não terá incentivos em revelá-las de outra maneira.

O que é importante ressaltar aqui, é que um *incentive compatible direct mechanism* não impõe restrições a nenhum tipo de jogo, ele apenas afirma que existem equilíbrios equivalentes para qualquer tipo de jogo. Dessa maneira, qualquer tipo de processo de barganha a ser analisado não estará limitado a uma forma específica de jogo, podendo ser analisado de maneira geral.

Para a formalização matemática, se tomará o exemplo mais comumente utilizado para analisar situações dentro da Teoria dos Conflitos Internacionais, uma disputa entre dois países (conflito bilateral) que pode levá-los a guerrear.

CAPÍTULO 3: MODELO *GAME-FREE* PARA CONFLITOS INTERNACIONAIS

3.1) Resultados e Formulações Iniciais

O modelo de Fey e Ramsay (2008, 2009) analisa o conflito como uma disputa entre países em que todas as variáveis - como território ou alocação de recursos - são divisíveis unitariamente, ou seja, são medidos a partir de unidades.

Em relação aos *payoffs* das nações em questão, elas dependerão das seguintes variáveis: probabilidade do país vencer o conflito; a utilidade de se vencer ou perder; e as ineficiências ou “desvantagens” em entrar em combate. Em relação à utilidade destes países em vencer ou perder, ela será tomada como 1, quando houver vitória e 0, quando houver derrota. O custo será definido como $c_i \geq 0$, sendo c_i o custo do país i em entrar num conflito. A probabilidade do país i vencer o conflito será representada por p_i . A utilidade do país i será função da probabilidade em vencer e dos custos em participar deste conflito, sendo representados por $w_i = p_i - c_i$.

Os países, dentro desta análise, possuem conhecimento privado em relação à sua capacidade de entrar ou não em uma guerra. O Estado i pode, por exemplo, ter conhecimento sobre o valor relativo em entrar em uma disputa, representado pelo custo relativo de um conflito c_i , ou então sobre a capacidade de sua força militar, determinada pela probabilidade deste país vencer a guerra p_i , ou ter informações privadas de ambas as situações.

Esse conhecimento privado pode ser representado pelo conceito de Harsanyi (1967) já apresentado de *types*, onde a informação privada do país i será representada por $t_i \in T_i \subset \mathbb{R}$. A distribuição conjunta de suas informações privadas, ou seja, de seus *types* será dada pela função $F(t)$, representando os pares $t = (t_1, t_2) \in T = T_1 \times T_2$. Alguns desses pares podem ser correlacionados, não sendo possível sua ocorrência nessa distribuição. Dessa maneira, limitam-se os pares possíveis ao conjunto T_p , representado pela função $F(t)$, excluindo os correlacionados.

Desta maneira, se os *payoffs* dos países são representados pela diferença entre suas probabilidades de vitória e seus custos de se engajar num conflito, sendo estas variáveis as representações de suas informações privadas, como demonstrado, expressadas pelos *types*, então, o *payoff* do país i pode ser definido como $w_i(t)$, função de seus *types*.

Do outro lado da análise, os países podem, ao invés de guerrear, resolver suas pendências de maneira pacífica (mas nem por isso amigável), através de barganhas, negociações diretas, mediações por outras nações ou organismos internacionais, ameaças, etc.

Cabe agora definir, de maneira abstrata, quais são as possíveis formas para determinar um jogo que represente de maneira adequada as soluções para as divergências entre as nações

Independentemente do processo a ser utilizado dentro de uma disputa (ou a guerra ou alguma solução pacífica) um jogo pode adotar alguma forma G , onde o conjunto de ações destes países (para fins de denotação, país 1 e país 2) podem ser representados, respectivamente, por A_1 e A_2 e uma função retorno (função resultado) determinada por estas ações, expressa por $g(A_1, A_2)$, onde $a_1 \in A_1$ e $a_2 \in A_2$.

Esta denominação permite a elaboração de um jogo que define o conjunto de ações para cada jogador e como o desenrolar desse jogo, ou seja, a interação entre as ações dos indivíduos determinam sua “função resultado”.

O resultado desse tipo de jogo, denominado por Fey e Ramsay (2008) de Jogo de Barganha em Conflitos⁶ será ou o estabelecimento de um acordo pacífico entre as nações ou um processo de embate.

Assim sendo, pode-se decompor os resultados desse jogo (decompor a função resultado) em duas partes, sendo aquela caso ocorra o conflito expressada pela probabilidade de guerra $\pi^g(a)$, e em caso de um acordo pacífico, o valor que foi estabelecido no acordo para o país um, $v^g(a)$. Supõe-se aqui também que qualquer tipo de acordo é considerado eficiente, e o valor estabelecido pelo acordo ao país 2 é análogo ao do país 1, sendo $1 - v^g(a)$. De uma forma geral, o valor estabelecido pelo acordo ao país i será $v_i^g(a)$.

⁶ Adaptação do inglês “Crisis Bargaining Game”

Desse modo, pode-se dizer que o *payoff* do país i que adota ações a dentro do conjunto A que lhe é definido, é estabelecido por:

$$u_i(a, t) = \pi^g(a)w_i(t) + (1 - \pi^g(a))v_i^g(a) \quad (1)$$

Ou seja, o *payoff* do país i é dado pela sua probabilidade de entrar em um conflito vezes o *payoff* que se pode obter ingressando neste conflito, somado à probabilidade de se obter um acordo pacífico multiplicada pelo valor estabelecido pelo contrato, sendo adotadas as ações em a .

Além disso, as informações privadas dos países podem, dentro da formulação de um acordo, alterar seu comportamento, ou seja, permitir mudanças em suas estratégias.

Fey e Ramsay (2009) analisam a questão de estratégia do modo como Myerson (1979) apresentou o modelo relacionando-os aos *types* dos indivíduos. A explicação parte da constatação de que as ações dos países são determinadas pelos *types* dos indivíduos, ou seja, por suas informações privadas. Além disso, as ações que os indivíduos tomam definem sua estratégia. Dessa maneira, pode-se formalmente estabelecer que a estratégia do país i segue uma função $s_i: T_i \rightarrow A_i$. Então, o conjunto de estratégias para o país i será dado por S_i , e o par de estratégias por $(s_1, s_2) = s \in S = S_1 \times S_2$.

Tendo estas constatações, pode-se obter um equilíbrio (no sentido bayesiano, de que o indivíduo tem conhecimento apenas de seu próprio *type*), em relação às suas estratégias. Um par de estratégias s^* é um equilíbrio se cada *type* de cada país (expressado por suas ações) for a melhor resposta para as estratégias dos demais países.

Então, para cada equilíbrio s^* de alguma forma de jogo G , $U_i(t_i)$ será a utilidade esperada do equilíbrio do *type* i do jogador i . Assim sendo:

$$U_i(t_i) = \int_{T_j} u_i(s^*(t_i, t_j), t) dF(t_j|t_i) \quad (2)$$

Sendo que $i \neq j$ e $u_i(a, t)$ são determinados pela equação precedente (1).

Dado estes resultados, é necessário agora estabelecer algumas definições para que seja possível formalizar o *revelation principle*.

Primeiramente, deve-se ter em mente uma forma de jogo que siga a estrutura analisada acima, fixando-se um equilíbrio determinado por s^* . Estabelecendo o conceito de que os *types* representam as estratégias dos países e suas ações representam suas estratégias, um equilíbrio s^* gera para um determinado par de *types* (t_1, t_2) um valor de equilíbrio de probabilidade de guerra $\pi(t) = \pi^g(s^*(t))$ e também um valor de equilíbrio para o valor estabelecido em um contrato igual a $v_i(t) = v_i^g(s^*(t))$.

Determinar as funções de probabilidade de guerra e de valor estabelecido em contrato nas formas $\pi(t)$ e $v(t)$ estabelece outro mecanismo, denominado *equivalent direct mechanism*, que representa que o espaço de ações dos indivíduos é apenas representado por seu *type*. Dessa maneira, se os agentes agirem apenas de acordo com seu *type*, ou seja, revelando de maneira correta suas informações pessoais, então este *direct mechanism* é também *incentive compatible*, no sentido já apresentado.

Vale lembrar que a utilização do *revelation principle* é fundamental para determinar equilíbrios que satisfaçam o *incentive-compatibility* em *direct mechanisms* e a partir de então estabelecer considerações gerais sobre os equilíbrios em todos os possíveis jogos de barganha entre países em conflito.

Para analisar essa questão, Myerson (1979) afirma que se um determinado equilíbrio (aqui denominado s^*) é um equilíbrio no sentido bayesiano dentro de alguma forma de jogo, então existe um *incentive-compatible direct mechanism* que estabelece o mesmo resultado desde jogo.

Assim sendo, pode-se estabelecer alguns conceitos necessários para a aplicabilidade do *revelation principle*.

Resultado 1: Se s^* é um equilíbrio no sentido bayesiano de algum determinado jogo de barganha G , então existe um *incentive-compatible direct mechanism* que promove o mesmo resultado.

Para explicar esse resultado, primeiramente, baseando-se num determinado equilíbrio (bayesiano) s^* dentro de uma determinada forma de jogo G , procura-se identificar um mecanismo que leve ao mesmo resultado deste equilíbrio, mas, fazendo com que os

agentes (países) baseiem suas estratégias em relação a seus *types*, para todo t . Isso é possível utilizando-se de um *direct-mechanism* que determine que os países exponham de maneira verdadeira suas informações privadas. Pode-se assim dizer que a função retorno desse jogo seja tal que $g(s^*(t))$, para todo $t \in T$.

Assim exposto, os países não terão incentivos a agir de maneira diferente, mentindo a respeito de suas informações. Isso porque se o país i mentir, dado seu *type* t_i dentro de um *direct mechanism*, a mentira não produzirá o mesmo resultado que ele poderá obter utilizando uma estratégia que corrobore o *direct mechanism*, pois fugiria do equilíbrio. Se assim o país bem entender, ele poderia mentir em qualquer formato que o jogo possa vir a tomar, não havendo a necessidade de se estabelecer um *direct mechanism*.

Para se implementar de maneira adequada estes princípios, é necessário primeiro avaliar algumas peculiaridades a respeito de barganhas entre países em conflito.

Fearon (1995) trabalhava a idéia de que os acordos deveriam ser voluntários entre os países, devido aos Estados não terem poder nem qualquer tipo de relações legislativas entre eles. Slantchev (2003) salientava o caráter anárquico das questões ligadas às relações internacionais, o que impedia determinados poderes coercitivos de um país a outro.

Fey e Ramsay (2009) prosseguem o modelo admitindo esse caráter. Os autores salientam que qualquer país pode, se bem entender, recusar um valor proposto por um determinado contrato $v^g(a)$ caso pense que poderá obter um retorno maior entrando em conflito.

Isso se dará caso a relação entre os retornos obtidos em guerrear ou estabelecer um acordo pacífico permita que algum dos lados prefira se engajar em um conflito. Em outras palavras, não há nenhuma forma de barganha, estabelecida com base em acordos voluntários que force algum país a estabelecer um acordo caso o mesmo visualize melhores frutos dentro de uma guerra. Formalmente, tem-se que um jogo de barganha tem acordos voluntários se, dado que $i \neq j$, existe uma determinada ação $\tilde{a}_i \in A_i$ tal que $v_i^g(\tilde{a}_i, a_j) = 1$ para todo $a_j \in A_j$, ou então, para todo $t \in T_p$, $v_i^g(\tilde{a}_i, a_j) \geq w_i(t)$, ou seja, o valor do acordo é maior ou igual ao retorno da guerra.

Outro ponto destacado pelos autores é que o processo de barganha dentro dos conflitos internacionais tem uma capacidade elevada de revelar as informações privadas das

nações enquanto é estabelecido. Essas informações são incorporadas pelos países, servindo tanto como catalisador na formação de um acordo como na rejeição do mesmo, optando mais veementemente pelo embate.

Pode-se dizer que, ao obter mais informações sobre seu(s) adversário(s), os países atualizam suas informações privadas. Esse processo pode ser representado por $u_i(v_i, t_i)$, sendo u_i a atualização das informações do país i a respeito das informações do país j , ou seja, de seu *type*, após observar o valor v_i estabelecido pelo acordo.

Agora, cabe inserir as constatações a respeito da peculiaridade dos acordos internacionais, que devem ser voluntários, dentro da Teoria do Desenho de Mecanismos, e identificar o *incentive-compatible direct mechanism* para a obtenção de equilíbrios pacíficos.

Outro resultado pode ser analisado:

Resultado 2: Suponha que s^* é um equilíbrio de uma forma qualquer de jogo de barganha em tempos de crise que tenha acordos voluntários. Logo, existe um *incentive-compatible direct mechanism* tal que $v_i(t) \geq E[w_i(t)|u_i(v_i, t_i)]$ para todo $t \in T_p$ tal que $\pi(t) \neq 1$.

Primeiramente, através da fórmula do *payoff* do país i , indicada pela equação (1), se a probabilidade de guerra $\pi^g(s^*(t)) \neq 1$, o valor de equilíbrio estipulado no contrato deve ser no mínimo igual ao *payoff* de ir à guerra.

Assim sendo, ao receber uma proposta de acordo igual a $v_i(t)$, o país i atualiza suas informações em relação ao seu adversário. Então, se mesmo assim, preferir ir à guerra (ação $\tilde{a}_i$), receberá um *payoff* esperado relativo a $w_i(t)$, dada sua atualização, igual a $E[w_i(t)|u_i(v_i, t_i)]$.

Então, para se obter um equilíbrio s^* onde não haja um desvio de ação (representado por $\tilde{a}_i$), dado que $\pi(t) = \pi^g(s^*(t)) \neq 1$, essa ação não pode ser vantajosa, ou seja, o retorno de um acordo pacífico deve ser pelo menos o mesmo do retorno da guerra, sendo possível apenas se $v_i(t) \geq E[w_i(t)|u_i(v_i, t_i)]$. Conclui-se com isto que, quando os acordos são apenas realizáveis de maneira voluntária, os países sempre vão ter a opção de atacar seus oponentes quando bem entenderem. Dessa maneira, o *payoff* do

acordo deve ser igual ou maior ao do conflito, incorporadas as informações adquiridas durante o processo de negociação.

Outro ponto que ainda necessita de análise é o fato de que, os Estados sempre terão a guerra como custosa (Schelling, 1960). Então, devem-se avaliar os casos em que as informações privadas sempre evitem os conflitos, e refutar aquele para os quais esse efeito não é visto. Para tanto, um dado equilíbrio s^* de um jogo de barganha em tempos de crise é sempre pacífico (evita a possibilidade de conflitos) sempre que $\pi^g(s^*(t)) = 0$, para todo $t \in T_p$.

Assim sendo, não há nenhum possível par de *types* que corrobore alguma ação que desvie do propósito de se obter uma resolução pacífica para um embate entre os países.

Deste modo:

Resultado 3: Dado um equilíbrio sempre pacífico s^* de um determinado jogo que possua acordos voluntários, há sempre um *incentive-compatible direct mechanism* para todo $t \in T_p$, $\pi(t) = 0$, e $v_i(t) \geq E[w_i(t)|u_i(v_i, t_i)]$.

Utilizando-se os pressupostos avaliados, agora é possível identificar resultados mais gerais para qualquer tipo de jogo (que represente uma barganha entre países em eminência de guerra) que possuam acordos voluntários. Para tanto se torna necessário identificar alguns conceitos que representem as estruturas das informações entre os países dentro do processo de negociação.

Para concluir esse capítulo, cabe apresentar outro resultado, expresso na forma de um lema, que será de suma importância para a análise matemática proposta do capítulo seguinte.

Lema 1: Suponha que os *types* dos países são correlacionados e que $T_i = [\underline{t}_i, \bar{t}_i]$. Além disso, suponha que $w_i(t)$ é continuamente diferenciável em t_1 e t_2 . Suponha também uma forma de jogo de barganha em tempos de crise G e que seu equilíbrio seja igual a s^* . Considerando $\pi(t) = \pi(s^*(t))$, tem-se que:

$$\frac{dU_i(t_i)}{dt_i} = \int_{T_j} \pi(t) \frac{\partial w_i}{\partial t_i}(t) dF_j(t_j).$$

Para provar este lema, supondo que s^* é o equilíbrio de uma forma de jogo qualquer G , existe, de acordo com o Resultado 1, um *incentive-compatible direct mechanism* que promove o mesmo equilíbrio s^* . Esse *direct mechanism* é dado por $\pi(t) = \pi^g(s^*(t))$ e $v_i(t) = v_i^g(s^*(t))$.

A utilidade esperada pelo *type* i pode ser expressa, de acordo com a equação 2 como:

$$U_i(t_i) = \int_{T_j} \pi(t_i, t_j) w_i(t_i, t_j) + (1 - \pi(t_i, t_j)) v_i(t_i, t_j) dF_j(t_j),$$

sendo $U_i(t_i) = U_i(t_i|t_i)$

Do mesmo modo, se o país revelar de maneira falsa um determinado *type* $\bar{t}_i$, a utilidade desse *type* será:

$$U_i(\bar{t}_i|t_i) = \int_{T_j} \pi(\bar{t}_i, t_j) w_i(t_i, t_j) + (1 - \pi(\bar{t}_i, t_j)) v_i(\bar{t}_i, t_j) dF_j(t_j).$$

Outro ponto é o de que, de acordo com a *incentive compatibility*, revelar de maneira inadequada suas informações dentro de um *direct mechanism* promove retornos no máximo iguais ao de revelar de maneira verídica sua informação privada. Assim:

$$U_i(t_i|t_i) \geq U_i(\tilde{t}_i|t_i),$$

e também

$$U_i(t_i) = U_i(t_i|t_i) = \max_{\tilde{t}_i \in T_i} U_i(\tilde{t}_i|t_i).$$

Aplicando-se agora o Teorema do Envelope (Milgrom e Segal, 2002), obtém-se:

$$\frac{dU_i(t_i)}{dt_i} = \left. \frac{dU_i(\tilde{t}_i|t_i)}{d\tilde{t}_i} \right|_{\tilde{t}_i=t_i} = \int_{T_j} \pi(\tilde{t}_i|t_j) \frac{\partial w_i}{\partial t_i}(t_i|t_j) dF_j(t_j) \Big|_{\tilde{t}_i=t_i}.$$

Ao igualarmos os valores de $\tilde{t}_i$ com t_i , é possível chegar ao resultado deste Lema.

3.2) Análise das Incertezas

Fey e Ramsay (2007) afirmam que a literatura dos conflitos internacionais sempre trouxe a incerteza como causa central de conflitos entre países, entretanto nunca definiu bem *que* tipos de incerteza são os mais fundamentais para se entrar ou não em guerra.

Como exemplo, Blainey (1988) estabelece como principal causa para um embate a incerteza quanto à força relativa entre os países. Por outro lado Schultz (1988) e Ramsay (2004) avaliam como sendo a incerteza a respeito dos custos da guerra o principal fator para um conflito.

A análise mais atual de Fey e Ramsay (2009) analisa a questão representando 2 tipos diferentes de incerteza. Entre elas, primeiramente, em relação aos custos de guerra e depois em relação à distribuição de poder entre as nações. Essas diferenças de incerteza promovem diferentes estruturas de informação dentro de um jogo, e dependendo de cada uma delas, diferentes estratégias por parte dos países no durante a realização do processo de barganha.

A incerteza em relação aos custos representa que os valores que os países têm em entrar em guerra não dependem dos mesmos valores de seus inimigos, ou seja, seus valores são independentes, não interferindo uns nos outros. Assim sendo, a preferência de um país entrar em guerra não afeta a distribuição dos *payoffs* dos adversários.

No caso em que a incerteza se dá na força relativa, diferentemente da incerteza em relação aos custos, os valores dos países em entrarem em guerra são interdependentes (permanecendo suas informações não correlacionadas).

3.2.1) Incerteza em Relação aos Custos da Guerra

Neste caso, as nações são incertas quanto ao custo de guerrear de seus inimigos. A questão também pode ser analisada como a incerteza quanto à “preferência por entrar num conflito” de um rival, ou seja, incerteza de quanto um país se sujeitaria para obter o bem, qualquer que seja o bem em disputa. (Schultz, 1988)

Vale ressaltar que as informações privadas dos indivíduos não são interdependentes, não influenciando as informações dos outros participantes. Em outras palavras, o conhecimento a respeito do custo de um determinado país não influencia diretamente o outro lado.

Yahri-Milo (2009) estudou um exemplo típico dessa incerteza, ocorrido durante os anos 30. O governo britânico estava incerto quanto ao poderio militar da Alemanha e do resto da Europa, e não tinha certeza quanto às intenções do Fuhrer e o quão longe ele estava disposto a levar um conflito.

Suponha que a probabilidade de um país, denominado país 1, seja igual a p , e seja conhecida por todos os jogadores. Entretanto, os custos de guerra dos países, c_i são desconhecidos de seus inimigos. Neste caso, suponha também que o *type* do país i é denominado pelo seu custo de guerra $c \in [0, \bar{c}_i] = C_i$, distribuído através de uma função $F_i(c_i)$, que representa C_i , e que os pares de *type* sejam $(c_1, c_2) = c \in C = C_1 \times C_2$.

Assim sendo, a probabilidade esperada de guerra para o país i , dada sua estratégia i é dada por:

$$\Pi(c_i) = E\pi^g(s(c_i, c_j)) dF(c_j) = \int_{C_j} \pi^g(s(c_i, c_j)) dF_j(c_j).$$

De maneira análoga, pode-se supor que a utilidade esperada de guerra, dada a mesma estratégia s , é

$$U_i(t_i) = \int_{c_j} u_i(s(c_i, c_j)c_i) dF_j(c_j).$$

Pode-se então, estabelecer o seguinte teorema:

Teorema 1: Suponha que os custos c_i correspondem às informações privadas dos indivíduos, e a probabilidade de vitória de cada país é conhecimento comum de todos. Suponha também que G é uma forma de jogo de barganha em tempos de crise qualquer e s^* é seu equilíbrio. Então $\Pi(c_i)$ e $U_i(c_i)$ são ambos decrescentes em c_i , para $i = 1, 2$.

Para expressá-lo, toma-se s^* como equilíbrio de qualquer forma de jogo de barganha em tempos de crise. De acordo com o Resultado 1 obtido anteriormente, existe um *incentive-compatible direct mechanism* que obtém o mesmo retorno do equilíbrio s^* proposto.

Já foi determinado também que esse mecanismo é definido por $\pi(c) = \pi^g(s^*(c))$ e $v_i(c) = v_i^g(s^*(c))$

De acordo com o resultado 2, a utilidade esperada to *type* c_i pode ser descrita como:

$$U_i(c_i) = \int_{c_j} \pi(c_i, c_j)(p_i - c_i) + (1 - \pi(c_i, c_j))v_i(c_i, c_j) dF_j(c_j)$$

Da mesma maneira, a utilidade de agir de maneira desonesta, através de um *type* $\tilde{c}_i$ pode ser descrita como:

$$U_i(\tilde{c}_i|c_i) = \int_{c_j} \pi(\tilde{c}_i, c_j)(p_i - c_i) + (1 - \pi(\tilde{c}_i, c_j))v_i(\tilde{c}_i, c_j) dF_j(c_j)$$

Além disso, o *incentive compatibility* ainda presume que para todo $c_i, \tilde{c}_i \in C_i$:

$$U_i(c_i|c_i) \geq U_i(\tilde{c}_i|c_i)$$

$$U_i(\tilde{c}_i|\tilde{c}_i) \geq U_i(c_i|\tilde{c}_i)$$

Adicionando as inequações às funções de utilidade e simplificando, temos que:

$$\begin{aligned} & \int_{c_j} \pi(c_i, c_j)(p_i - c_i) - \int_{c_j} \pi(\tilde{c}_i, c_j)(p_i - \tilde{c}_i) \\ & \geq \int_{c_j} \pi(\tilde{c}_i, c_j)(p_i - c_i) - \int_{c_j} \pi(\tilde{c}_i, c_j)(p_i - \tilde{c}_i) \\ & \int_{c_j} \pi(c_i, c_j)(\tilde{c}_i - c_i)dF_j(c_j) \geq \int_{c_j} \pi(\tilde{c}_i, c_j)(\tilde{c}_i - c_i)dF_j(c_j) \end{aligned}$$

Dessa maneira, se $\tilde{c}_i \geq c_i$, então $\Pi(\tilde{c}_i) \geq \Pi(c_i)$, provando que $\Pi(c_i)$ é decrescente.

Para $U_i(c_i)$, aplicando-se o Lema 1, obtém-se:

$$\frac{dU_i(c_i)}{dc_i} = \int_{c_j} \pi(c)(-1)dF_j(c_j).$$

Como $\pi(c)$ é sempre não-negativo (pois é uma probabilidade), então, elevações nos valores de c promovem valores negativos da utilidade, devido ao multiplicador (-1) da função, comprovando-se que $U_i(c_i)$ é também decrescente.

Assim sendo, Fey e Ramsay (2009) salientam como uma das características fundamentais de modelos de jogos de barganha em tempos de crise a monotonicidade

das funções de utilidade e probabilidade de guerra. Desse modo, os resultados demonstram que em qualquer equilíbrio de qualquer forma de jogo que represente disputas entre países, variações em seus *types* (ou seja, em suas informações privadas) levam a mudanças em suas funções de utilidade, e assim variações em seus *payoffs*.

Considerando-se agora equilíbrios pacíficos, onde a probabilidade de guerra é igual a zero, cabem algumas considerações a respeito de suas características.

Teorema 2: Se os custos c_i correspondem às informações privadas dos indivíduos, e a probabilidade de vitória de cada país é conhecimento comum de todos, então existe uma forma jogo de barganha em tempos de crise que possui *acordos voluntários* onde sempre haverá *equilíbrios pacíficos*.

Toma-se como exemplo, uma forma de jogo que tem por definição acordos voluntários e que sempre viabilizará equilíbrios pacíficos.

Considere uma forma de jogo que tenha acordos voluntários tais que $v_i^g(a) = p_i$ para todo $a \in A$ e $\pi^g(a *) = 0$ para alguma ação $a * \in A$. Assim sendo, $s * (c) = a *$ é um equilíbrio, pois se qualquer um dos países tentar mudar de estratégia e iniciar uma guerra, ambos sairão perdendo, e se tentarem se deslocar para outro equilíbrio, os *payoffs* continuarão os mesmos. Esse equilíbrio é, por estrutura, pacífico.

Para uma formulação mais geral, ou seja, para representar equilíbrios pacíficos em todos os possíveis jogos que possuam acordos voluntários, pode-se estabelecer que para o país i , o valor esperado de um acordo, dada uma ação a_i e uma estratégia empregada por seu adversário s_j , pode ser descrito como:

$$Ev_i^g(a_i | s_j) = \int_{c_j} v_i^g(a_i s_j(c_j)) dF_j(c_j).$$

O teorema seguinte permite obter as condições necessárias para que sempre haja equilíbrios pacíficos:

Teorema 3: Assumem-se custos c_i correspondem às informações privadas dos indivíduos, e a probabilidade de vitória de cada país é conhecimento comum de todos.

Além disso, considera-se G uma forma de jogo de barganha em tempos de crise qualquer, onde há acordos voluntários. Dessa maneira, um equilíbrio s^* de G é sempre pacífico apenas se, para $i = 1, 2$; $E v_i^g(s_i^*(c_i)|s_j^*) = p_i$ para todo $c_i \in C_i$; $E v_i^g(a_i|s_j^*) \leq p_i$ para todo $a_i \in A_i$ e $v_i^g(s^*(c)) \geq p_i - c_i$ para todo $c \in C$.

Para prová-lo, suponha que a mesma estratégia s^* é um equilíbrio pacífico de uma forma de jogo qualquer (jogo de barganha) que tenha acordos voluntários entre os países. Como já definido no Resultado 3, há um *incentive-compatible direct mechanism* tal qual $\pi(c) = 0$ e

$$v_i(c) \geq E[w_i(c)|\mu_i(v_i, c_i)] = p_i - c_i,$$

sendo esta condição para todo $c_i \in C_i$ e para $i = 1, 2$.

Em relação à $E[w_i(c)|\mu_i(v_i, c_i)] = p_i - c_i$, os valores são iguais, pois os valores de se entrar em uma guerra do país i não dependem dos valores do país j , representados por c_j . O *direct mechanism* neste caso é dado por $v_i(c) = v_i^g(s^*(c))$ e $\pi(c) = \pi^g(s^*(c))$.

Agora analisando a utilidade, uma questão importante é a de que, como o $\pi(c) = 0$, ao aplicar o Lema 1, a derivada da utilidade também resulta em zero ($U_i'(c_i) = 0$), portanto a utilidade $U_i(t_i) = \int_{C_j} u_i(c_i, c_j) dF_j(c_j)$ é constante em c_i .

Outro ponto fundamental é o de que $v_i(c) \geq p_i - c_i$ deve valer para $c_i = 0$. Então

$$\int_{C_j} u_i(c_i, c_j) dF_j(c_j) = \int_{C_j} u_i(0, c_j) dF_j(c_j) \geq p_i,$$

para todo $c_i \in C_i$ e $i = 1, 2$.

Para verificar se os valores de v são iguais a 1, suponha que eles não sejam. Assim:

$$\begin{aligned}
& \int_{c_1} \int_{c_2} [v_1(c) + v_2(c)] dF_2(c_2) dF_1(c_1) \\
&= \int_{c_1} \int_{c_2} v_1(c) dF_2(c_2) dF_1(c_1) + \int_{c_1} \int_{c_2} v_2(c) dF_1(c_1) dF_2(c_2) \\
&> \int_{c_2} p dF_2(c_2) + \int_{c_1} (1-p) dF_1(c_1) = p + (1-p) = 1.
\end{aligned}$$

Como todo acordo pacífico é eficiente, $v_1(c) + v_2(c)$ devem ser igual a 1, para todo $c \in C$, o que não ocorre aqui, provando assim o Teorema 3 por contradição.

Esse teorema revela algumas proposições importantes em relação a equilíbrios pacíficos. Primeiramente, equilíbrios pacíficos devem gerar estruturas de *payoff* completamente insensíveis em relação às informações privadas dos países, ou seja, em relação aos seus custos de guerra. Em outras palavras, independentemente se os custos dos países são altos ou baixos, devem receber as mesmas propostas de acordo, corroborando os mesmos valores.

Assim sendo, equilíbrios pacíficos dependem apenas das informações que estão disponíveis publicamente e não podem variar em relação às informações privadas.

Entretanto, o que acontece é que muitas vezes, não é tão simples encontrar *payoffs* iguais para *types* altos e baixos (ou seja, *types* que apresentam custos de guerrear altos e baixos). O que se quer dizer aqui é que algumas vezes, em conflitos reais, incertezas em relação aos custos podem levar os países ao combate. No exemplo acima, considerou-se probabilidade de guerra igual a zero ($\pi(c) = 0$). Para reforçar o teorema, podemos supor a seguinte condição:

Corolário 1: Assumem-se custos c_i correspondentes às informações privadas dos indivíduos, e a probabilidade de vitória de cada país é conhecimento comum de todos. Além disso, suponha que G é uma forma de jogo de barganha em tempos de crise qualquer, onde há acordos voluntários. Além disso, considere s^* um equilíbrio qualquer de G . Então $U_i(c_i) \neq U_i(\bar{c}_i)$ para algum $c_i, \bar{c}_i \in C_i$ se e somente se existe probabilidade de guerra positiva no equilíbrio.

Esse corolário supõe que haja probabilidade de guerra. Assim sendo, a probabilidade de guerra do modelo pode ser representada por $\pi^g(s^*(c)) \geq 0$. Aplicando-se o Lema 1 do mesmo modo como foi aplicado para o Teorema 3, pode-se observar que para os casos em que a probabilidade de guerra é positiva, tem-se $U'_i(c_i) > 0$.

Em outras palavras, pode-se dizer que este corolário estabelece que só nos casos em que a probabilidade de guerra é positiva é que dois *types* diferentes (como no exemplo, um alto e um baixo) poderão proporcionar *payoffs* diferentes. Se os países esperam melhores condições indo à guerra, assim o farão.

3.2.2) Incerteza em Relação à Força Relativa

Nesta seção, será descrito o modelo de Mark e Fey (2008, 2009) para outro tipo de incerteza comumente debatido dentro da teoria dos conflitos internacionais. Na seção anterior, foram apresentados pressupostos gerais para os casos em que a incerteza se dava em relação aos custos em guerrear. Agora, serão apresentados os conceitos em relação à distribuição de força entre os países e a probabilidades dos mesmos de vencer o conflito.

Primeiramente, cabe constatar que os agentes, dentro desse mecanismo, têm conhecimento sobre os custos relativos de guerra de seus adversários, mas são incertos quanto aos resultados desses possíveis conflitos. Assim sendo, os países possuem informações privadas em relação a, por exemplo, sua força militar (tecnologia, preparação de contingente, etc.) ou em relação à sua estratégia de combate. Essas informações levam os países a manter secretas suas crenças em relação aos resultados de entrar em guerra e conseqüentemente, a sua probabilidade de vencer. A incerteza em relação à força relativa leva a uma situação em que os valores dos *payoffs* são interdependentes e os *types* são não-correlacionados. Desse modo, os valores de determinado país não são dependentes apenas de seu *type*, mas também dos *types* das outras nações. Como os *types* não são correlacionados, o *type* de um determinado país não influencia a determinação do *type* de outro qualquer.

O modelo de Fey e Ramsay (2009) para incerteza em relação à força relativa se inicia assumindo que os custos de se guerrear de dois países (1 e 2) são dados por c_1 e c_2 , e

são de conhecimento comum, mas os países possuem informações privadas em relação à sua força militar, que resulta em informações privadas quanto à probabilidade de vencer a guerra.

Para implementar o modelo, supõe-se que o *type* do país i é dado por $t_i \in [\underline{t}_i, \bar{t}_i] = T_i$, distribuído de acordo com uma função F_i e a probabilidade do país 1 um de vencer a guerra é dada por $p(t_1, t_2)$, ou seja, ela depende não apenas de seu próprio *type*, mas também do de seu adversário.

Além disso, pode-se dizer que as probabilidades de vencer a guerra dos países são monotônicas em relação a seus *types*, ou seja, quanto maior for o *type* de determinado país, maior será sua probabilidade de vencer a guerra, sendo tudo o mais constante.

Formalmente, pode-se expressar que sendo $t_1 > t'_1$, então $p(t_1, t_2) \geq p(t'_1, t_2)$, para todo $t_2 \in T_2$.

Se maiores *types* do país 1 levam a maiores probabilidades de vencer a guerra, (probabilidade monotônica em relação aos *types*) maiores *types* de seu adversário levam à diminuição de sua probabilidade de prevalecer num conflito, ou seja, ela é monotonicamente decrescente em relação ao *type* do país 2.

Outro ponto importante é em relação à atualização das informações de um país em relação aos outros, como já foi demonstrado. Um país, ao receber uma oferta de acordo, pode atualizar-se em relação às informações privadas de seu adversário, inferindo em relação às preferências demonstradas pelos mesmos em sua proposta.

Como já visto, $\mu_i(v_i, t_i)$ expressa a atualização do país i em relação ao *type* do país j , após este país fazer uma proposta de acordo v_i . Toma-se aqui $V_1(t_1, v) = \{t_2 | (t_1, t_2) = v\}$ como o conjunto de todos os possíveis *types* do país 2 que o país 1 pode inferir ser possível após analisar a proposta de acordo v feita pelo país 2. Assim sendo, a utilidade esperada pelo país 1 em participar de um conflito pode ser expressa como:

$$E[w_1(t)|\mu_1(v_1, t_1)] = E[p(t_1, t_2)|V_1(t_1, v_1)] - c_1. \quad (3)$$

O lado direito expressa a probabilidade do país 1 ganhar a guerra (dependente dos *types* de ambos os países) dado o conjunto dos possíveis *types* inferidos pelo mesmo país 1, subtraídos pelo custo. De forma análoga, o mesmo resultado pode ser obtido para o país 2.

Pode-se estabelecer ainda que:

$$P_1(t_1) = \int_{T_2} p(t_1, y) dF_2(y)$$

Sendo a probabilidade do país 1 em vencer o conflito, dado seu *type* t_1 e

$$P_2(t_2) = \int_{T_1} p(x, t_2) dF_1(x),$$

a probabilidade do país 2 em perder a guerra, dado seu *type* t_2 .

Subtraindo essas duas probabilidades, para determinados *types* $\bar{t}_1$ e $\bar{t}_2$, pode-se determinar que, quanto menor o valor dessa diferença, maior a possibilidade de haver conflitos. Em outras palavras, maior a distância entre a vitória do país 1 e a derrota do país 2. O mesmo pode-se dizer em relação ao país 2. Em termos de custo, pode-se dizer que:

$$\bar{c} = P_1(\bar{t}_1) - P_2(\bar{t}_2)$$

Então, quanto menor o valor da subtração, menor será o custo de entrar em conflito. O Teorema 4 demonstra que se os custos dos países forem menores do que o custo $\bar{c}$, então sempre haverá probabilidade positiva de guerra, não importa o mecanismo imposto.

Teorema 4: Suponha que os custos são de conhecimento comum a todos os países, mas cada país é incerto em relação à probabilidade de vencer a guerra. Se $c_1 + c_2 < \bar{c}$, então não haverá sempre um equilíbrio pacífico em qualquer forma de jogo de barganha em tempos de crise que tenha acordos pacíficos.

Para provar o teorema, por um processo de contradição, supomos que sempre haverá equilíbrios pacíficos dentro de uma forma de jogo de barganha em tempos de crise que tenha acordos voluntários.

Toma-se aqui, pelo Resultado 3, que existe um *incentive-compatible direct mechanism* que corrobore uma probabilidade de guerra igual a zero ($\pi(t) = 0$) e $v_i(t) \geq E[w_i(t)|\mu_i(v_i, t_i)]$ para todo t e $i = 1, 2$. Dessa maneira, a utilidade esperada do país 1, sendo seu *type* igual a t_1 , mas revelando contraditoriamente um *type* igual a $\hat{t}_1$, dada a probabilidade de guerra igual a zero:

$$U_1(\hat{t}_1|t_1) = \int_{T_2} p(\hat{t}_1, y) dF_2(y)$$

Como já observado como condição da *incentive compatibility*, tem-se:

$$U_i(t_i|t_i) \geq U_i(\hat{t}_i|t_i). \quad (4)$$

Para todo $t_1, \hat{t}_1 \in T_1$.

Como se pode observar pela equação 4, a função-utilidade não depende do *type* t_1 do país. Assim, sendo o único modo de satisfazer a condição necessária da *incentive compatibility* é se a utilidade $U_i(\hat{t}_i|t_i)$ for sempre constante, para todo $\hat{t}_i$ e t_i , não provocando mudanças na indiferença proposta. Para nomear essa constante, utiliza-se o símbolo $\bar{U}_1$.

Em relação à condição de que $v_i(t) \geq E[w_i(t)|\mu_i(v_i, t_i)]$ para todo t e $i = 1, 2$, tem-se, da equação 3 a igualdade com $E[p(\bar{t}_1|t_2)|\mu_i(v_i, t_i)] - c_1$. Substituindo os valores de t_1 por um *type* igual a $\bar{t}_1$, tem-se:

$$v_i(\bar{t}_1|t_2) \geq E[p(\bar{t}_1|t_2)|\mu_i(v_i, t_i)] - c_1.$$

Adicionando-se a esperança aos dois lados da equação:

$$E[v_i(\bar{t}_1|t_2)] \geq E[E[p(\bar{t}_1|t_2)|\mu_i(v_i, t_i)]] - c_1.$$

Através da lei das expectativas iteradas, pode se dizer que a equação 3 equivale a:

$$\bar{U}_1 = \int_{T_2} v_1(\bar{t}_1, t_2) dF_2(t_2) \geq \int_{T_2} p(\bar{t}_1, t_2) dF_2(t_2) - c_1.$$

Da mesma maneira, se supusermos as mesmas condições para o país 2, denominando $t_2 = \bar{t}_2$, sua utilidade será:

$$\bar{U}_2 = \int_{T_1} v_2(t_1, \bar{t}_2) dF_1(t_1) \geq \int_{T_1} [1 - p(t_1, \bar{t}_2)] dF_1(t_1) - c_2$$

Agora, adicionando o fato de que o valor dos acordos estipulados entre os países somados devem ser igual a 1 ($v_1(t_1, t_2) + v_2(t_1, t_2) = 1$), pois trata-se de um jogo com acordos voluntários que sempre gera equilíbrios pacíficos, tem-se que:

$$\int_{T_1} \int_{T_2} [v_1(t) + v_2(t)] dF_2(t_2) dF_1(t_1) = 1$$

$$\int_{T_1} \int_{T_2} v_1(t) dF_2(t_2) dF_1(t_1) + \int_{T_2} \int_{T_1} v_2(t) dF_1(t_1) dF_2(t_2) = 1$$

$$\int_{T_1} \bar{U}_1 dF_1(t_1) + \int_{T_2} \bar{U}_2 dF_2(t_2) = 1$$

$$\bar{U}_1 + \bar{U}_2 = 1$$

Obteve-se aqui que a soma das utilidades dos indivíduos $\bar{U}_1 + \bar{U}_2$ também é igual a 1. Então, adicionando-se as inequações que são requisitos da *incentive-compatibility*, determina-se

$$1 \geq \int_{T_2} p(\bar{t}_1, y) dF_2(y) - c_1 + 1 - \int_{T_2} p(x, \bar{t}_2) dF_1(x) - c_2,$$

e então,

$$c_1 + c_2 \geq \bar{c}$$

Logo, em jogos onde ocorrem acordos voluntários que sempre geram equilíbrios pacíficos, com probabilidade de guerra igual a zero, os valores dos custos dos países 1 e 2 somados são $c_1 + c_2 \geq \bar{c}$. Por contradição, custos menores a $\bar{c}$ sempre gerarão probabilidade de guerra, provando o teorema 3.

Essas constatações determinam que se $c_1 + c_2 < \bar{c}$, então qualquer que seja a interação entre os países, sempre que houver informações privadas em relação à força relativa dos países, que infere em sua probabilidade de vitória, sempre haverá a probabilidade de guerra maior que zero.

Por um argumento semelhante, pôde-se observar que se $c_1 + c_2 \geq \bar{c}$, então é possível evitar que se ocorra um conflito. Logo, pelo Teorema 5:

Teorema 5: Supondo que os custos dos países são de conhecimento comum mas cada país é incerto em relação a sua probabilidade de vencer o conflito, se $c_1 + c_2 \geq \bar{c}$, então há uma forma de jogo de barganha em tempos de crise que tem acordos voluntários onde sempre existe um equilíbrio pacífico.

Do mesmo modo que o apresentado após o teorema 2, considere uma forma de jogo que tenha acordos voluntários tais que $v^g(a) = P_1(t_1) - c_1$ para todo $a \in A$ e $\pi^g(a^*) = 0$ para alguma ação $a^* \in A$. Assim sendo, $s^*(t) = a^*$ é um equilíbrio, pois se tentarem se deslocar para outro equilíbrio pacífico, os *payoffs* continuarão os mesmos. Esse equilíbrio é, novamente, por estrutura, pacífico.

Se por exemplo, o país 1 ao invés de buscar outra ação pacífica, decidir ir a guerra, dado seu *type* t_1 , considerando o *type* $\bar{t}_1$ referente a ações pacíficas, tem-se:

$$\int_{T_2} p(t_1, y) dF_2(y) - c_1 \leq \int_{T_2} p(\bar{t}_1, y) dF_2(y) - c_1 = P_1(\bar{t}_1) - c_1 = g_v(a^*)^7.$$

Dessa maneira, desvios na conduta pacífica (lado esquerdo da equação) sempre levaram a resultados no máximo iguais aos da estratégia pacífica. Assim sendo, os países (aqui representado pelo país 1) não terão incentivos para mudar de estratégia.

Assim sendo, pode-se estabelecer resumidamente quando é possível ou quando não é possível alcançar equilíbrios pacíficos. Se os custos forem maiores do que $\bar{c}$, então é possível alcançar esse tipo de equilíbrio; se forem menores, torna-se impossível.

Outro ponto importante é que o valor deste custo $\bar{c}$, que estabelece o limite de se entrar ou não em conflito, não depende de *nenhuma forma de jogo*, apenas das distribuições de F_1 , F_2 e de $p(t_1, t_2)$, dando o caráter geral do modelo.

Ainda outro ponto pode complementar a análise. Supondo o caso onde os custos dos países são menores do que $\bar{c}$, possibilitando o conflito. Pode haver uma terceira parte dentro desse embate, que esteja interessada num acordo pacífico e que pode gerar subsídios para que esse processo ocorra:

Corolário 2: Se uma terceira parte prover o um subsídio tal que:

⁷ Lembrando-se que g_v representa a função retorno, que é a soma do valor do acordo proposto e do retorno em se guerrar. Como o equilíbrio em questão é pacífico, logo $v^g = g_v$.

$$\phi \geq P_1(\bar{t}_1) - P_2(\bar{t}_2) - (c_1 - c_2),$$

então há uma forma de jogo de barganha em tempos de crise que tem acordos voluntários onde sempre existe um equilíbrio pacífico.

Schelling (1960) afirma que quanto mais poderosos são os países interessados em não haver conflito, maiores serão os subsídios e mais difícil torna-se a ocorrência de guerras.

Em síntese, pode-se estabelecer que:

Teorema 6: Suponha que os custos dos países são de conhecimento comum mas cada país é incerto em relação à sua probabilidade de vencer o conflito. Além disso, suponha que G é forma de jogo de barganha em tempos de crise qualquer, onde há acordos voluntários. Além disso, considere s^* um equilíbrio qualquer de G . Então $U_i(t_i) \neq U_i(\bar{t}_i)$ para algum $t_i, \bar{t}_i \in T_i$ se e somente se existe probabilidade de guerra positiva no equilíbrio.

Esse teorema demonstra resultados semelhantes aos resultados obtidos em relação a incerteza quanto aos custos. Assim sendo, a utilidade do país i só será diferente da obtida pelo equilíbrio pacífico se houver probabilidade de guerra positiva. Caso contrário, se os indivíduos não tiverem incentivos para mudar de estratégia, não o farão.

CONCLUSÕES

Um dos pontos fundamentais do modelo apresentado é que o método *game-free* empregado é capaz de gerar resultados pra qualquer tipo de interação entre países que se encontram em conflito eminente. Mais detalhadamente, o método tem aplicação em qualquer tipo de jogo de barganha onde os países possuam incerteza em relação aos valores de seus adversários e, dado este clima de instabilidade, podem ofertar e contra-ofertar propostas a fim de determinar acordos pacíficos, ou então entrar em combate. Também se torna possível aferir resultados quando uma terceira parte participa do processo de barganha, exercendo subsídios e incentivos para que o conflito não ocorra. Além das incertezas, o modelo permite analisar os incentivos que levam os países a representar de maneira desonesta suas informações, impedindo a formação de resoluções pacíficas.

Outro fator importante estabelecido foi a capacidade do modelo de estabelecer, racionalmente, a guerra como consequência de informações privadas e dos incentivos para revelar de maneira desonesta suas informações. Mais do que isso, foi possível também identificar que as informações privadas dos indivíduos são um empecilho para a formação de acordos pacíficos dependendo dos tipos de incerteza enfrentada pelos países.

O modelo de Fey e Ramsay possibilitou justificar de maneira geral as causas que determinam se os países participam de um confronto ou não, sem a necessidade de se apegar a modelos de jogos específicos.

Assim como demonstrado nos teoremas, os resultados dos *payoffs* e a consequente determinação de equilíbrios pacíficos não dependem de nenhuma especificidade do jogo. Como exemplo, no caso de incerteza em relação aos custos de guerra, de acordo com o teorema 5, a existência ou não de equilíbrios pacíficos não depende de nenhuma forma específica de jogo, mas apenas dos custos dos países em entrar em conflito e da incerteza em relação à forma relativa dos países, expressa, por exemplo, na forma militar, estratégica, etc.

Aplicações deste tipo de modelo em conflitos internacionais são bem atuais, elucidando o caráter de fronteira dessa literatura dentro das relações entre os países. Aqui foi

possível analisar de maneira geral como os países se relacionam com diferencial de informações e como são estabelecidos de maneira geral os resultados de suas ações.

Os exemplos apresentados aqui representaram *types* unidimensionais, ou seja, os países tinham incerteza e guardavam informações privadas em relação a apenas uma variável: ou em relação ao custo de guerra, ou à força relativa. Pode-se estabelecer também modelos que considerem *types* multidimensionais, onde os agentes possuam informações privadas a respeito tanto dos custos quanto da força relativa, ou até de novas variáveis a serem incorporadas.

Como observado por Fey e Ramsay (2008), esse tipo de abordagem ainda necessita de estudos, resultando assim em futuras pesquisas que possam enriquecer o debate.

BIBLIOGRAFIA

AUMANN, R.J. (1987) "Correlated Equilibrium as an Expression of Bayesian Rationality," *Econometrica*, Vol. 55, N°. 1, pp. 1-18

BANKS, J. S. (1990) "Equilibrium Behavior in Crisis Bargaining Games," *American Journal of Political Science* Vol. 34, n°.3, pp. 599-614.

BINMORE, K. (2009) "Interpersonal Comparison of Utility," *Centre for Economic Learning and Social Evolution*, pp. 1-20.

BLAINEY, G. (1988) "The Causes of War," NY: The Free Press.

DASGUPTA, P., HAMMOND, P., MASKIN, E. (1979) "The implementation of social choice rules: some general results on incentive compatibility," *Review of Economic Studies*, Vol. 46, pp. 185-216.

FEARON, J.D. (1995) "Rationalist explanations for war," *Int. Organ*, Vol.49, pp. 379–414.

FEY, M., RAMSAY, K. W. (2007) "Mutual Optimism and War," *American Journal of Political Science*, Vol.51, n°. 4, pp.:738-754.

FEY, M., RAMSAY, K. W. (2008) "Uncertainty and Incentives in Crisis Bargaining," Typescript, University of Rochester.

FEY, M., RAMSAY, K. W.(2009) "Mechanism design goes to war: peaceful outcomes with interdependent and correlated types," *Review of Economic Design*, Vol.3, n°. 3, pp. 233-250.

HARSANYI, J. (1967) "Games with Incomplete Information Played by 'Bayesian' Players, I-III. Part I. The Basic Model," *Management Science*, Vol. 14, n°. 3, pp. 159-182.

HURWICZ, L. (1972) "On informationally decentralized systems," in: MCGUIRE, C. B., RADNER, R., eds., *Decision and Organization: a Volume in Honor of Jacob Marshak*, North-Holland, pp. 297-336.

KREPS, D. (1994) "A course of microeconomics theory," Harvester Wheatsheaf.

- MILGRON, P., SEGAL, I. (2002) "Envelope Theorems for Arbitrary Choice Sets." *Econometrica*, Vol. 70, n°. 2, pp. 583-601.
- MITCHELL, R. B. (1994) "Regime design matters: intentional oil pollution and treaty compliance". *Int Organ*, Vol. 48, pp. 425–458.
- MYERSON, R. B. (1979) "Incentive Compatibility and the Bargaining Problem," *Econometrica*, Vol.47, pp. 61-74.
- MYERSON, R. B. (1982) "Optimal Coordination Mechanisms in Generalized Principal-Agent Problems," *Journal of Mathematical Economics*, Vol.10, pp. 67-81.
- MYERSON, R. B. (1988) "Mechanism Design," Center for Mathematical Studies in Economics and Management Science; Center for Mathematical Studies in Economics and Management Science, Northwestern University, Discussion paper n°. 796.
- MYERSON, R.B., SATTERHWAITE, M.A. (1983) "Efficient mechanisms for bilateral trading," *J Econ Theory*, Vol. 29, pp. 265–281.
- NOBEL (2007) "Mechanism Design Theory," Compiled by the Prize Committee of the Royal Swedish Academy of Sciences.
- PALFREY, T., SRIVASTAVA, S. (1987) "On Bayesian implementable allocations," *Review of Economic Studies*, Vol. 54, pp. 193-208.
- POWELL, R. (2004) "Bargaining and Learning While Fighting," *American Journal of Political Science* Vol. 48, n°.2, pp. 344-361.
- RAMSAY, K. W. (2004) "Politics at the Water's Edge: Crisis Bargaining and Electoral Competition," *Journal of Conflict Resolution*, Vol. 48, n°.4, pp.459-486.
- SCHELLING, T. C. (1960) "Strategy of Conflict," Harvard U Press.
- SCHULTZ, K. (1998) "Domestic Opposition and Signaling in International Crisis," *American Political Science Review*, Vol. 94, n° 4, pp. 829-844.
- SLANTCHEV, B. L. (2003) "The Principle of Convergence in Wartime Negotiation," *American Journal of Political Science*, Vol. 97, n°. 4, pp. 621-632.

WALTZ, K. (2001) "Man, the State and War: A Theoretical Analysis," Columbia University Press.

YARHI-MILO, K. (2009) "Knowing Thy Adversary: Assessment of Intentions in International Politics," PhD Dissertation, University of Pennsylvania.