

Editorial Manager(tm) for Computational and Mathematical Organization Theory
Manuscript Draft

Manuscript Number: CMOT32

Title: Modeling Factions for 'Effects Based Operations': Part II - Behavioral Game Theory

Article Type: Manuscript

Keywords: political simulation, agent-based models; behavioral game theory; validation; policy analysis tools

Corresponding Author: Barry Silverman,

Corresponding Author's Institution: Towne Rm 251

First Author: Barry Silverman

Order of Authors: Barry Silverman; Barry G Silverman; Gnana K Bharathy; Benjamin Nye; Tony E Smith

Modeling Factions for ‘Effects Based Operations’: Part II – Behavioral Game Theory

Barry G. Silverman¹, Gnana Bharathy¹, Benjamin Nye¹, Roy J. Eidelson²

1- Electrical and Systems Engineering Dept.

2- Asch Center for EthnoPolitical Conflict

University of Pennsylvania, Philadelphia, PA 19104-6315

basil@seas.upenn.edu

W: (215) 573-8368

January 2007

Modeling Factions for ‘Effects Based Operations’:

Part II – Behavioral Game Theory

“game theory is no better than chance (at predicting real world conflict)”
- JS Armstrong (2002), p. 345

Military, diplomatic, and intelligence analysts are increasingly interested in having a valid system of models that span the social sciences and interoperate so that one can determine the effects that may arise from alternative operations (courses of action) in different lands. Part I of this article concentrated on internal validity of the components of such a synthetic framework – a world diplomacy game as well as the agent architecture for modeling leaders and followers in different conflicts. But how valid are such model collections once they are integrated together and used out-of-sample (see Section 1)? Section 2 compares these realistic, descriptive agents to normative rational actor theory and offers equilibria insights for conflict games. Sections 3 and 4 offer two real world cases (Iraq and SE Asia) where the agent models are subjected to validity tests and an EBO experiment is then run for each case. We conclude by arguing that substantial effort on game realism, best-of-breed social science models, and agent validation efforts is essential if analytic experiments are to effectively explore conflicts and alternative ways to influence outcomes. Such efforts are likely to improve behavioral game theory as well.

Keywords: political simulation, agent-based models; game theory; validation; policy analysis tools

1) Introduction and Purpose

Analytic game theory is the mathematics of strategy, and as such, provides a language that is both rich and crisp. At the same time, analytic game theory has an abysmal record of explaining and/or predicting real world conflict – about the same as random chance according to Armstrong (2002), Green (2002). In the field of economics, Camerer (2003) points out that the explanatory and predictive powers of analytic game theory are being improved by replacing prescriptions from rational economics with descriptions from the psychology of monetary judgment and decision making. This has resulted in ‘behavioral game theory’ which adds in emotions, mistakes, heuristics, and so on. In this paper, we pursue the same approach and believe the term ‘behavioral game theory’ or just BGT is broad enough to cover all areas of social science, not just economics.

Specifically, the military, diplomatic, and intelligence analysis community would like for behavioral game theory to satisfy a wide and expanding range of scenario modeling and simulation concerns. Their interest goes beyond mission-oriented military behaviors, to also include simulations of the effects that an array of alternative diplomatic, intelligence, military, and economic (DIME) actions might have upon the political, military, economic, social, informational (psyops), and infrastructure (PMESII) dimensions of a foreign region. The goal is

to understand factional tensions and issues, and to examine alternative ways to influence and possibly shape outcomes for the collective good. There are at least four general uses for simulation of alternatives, including:

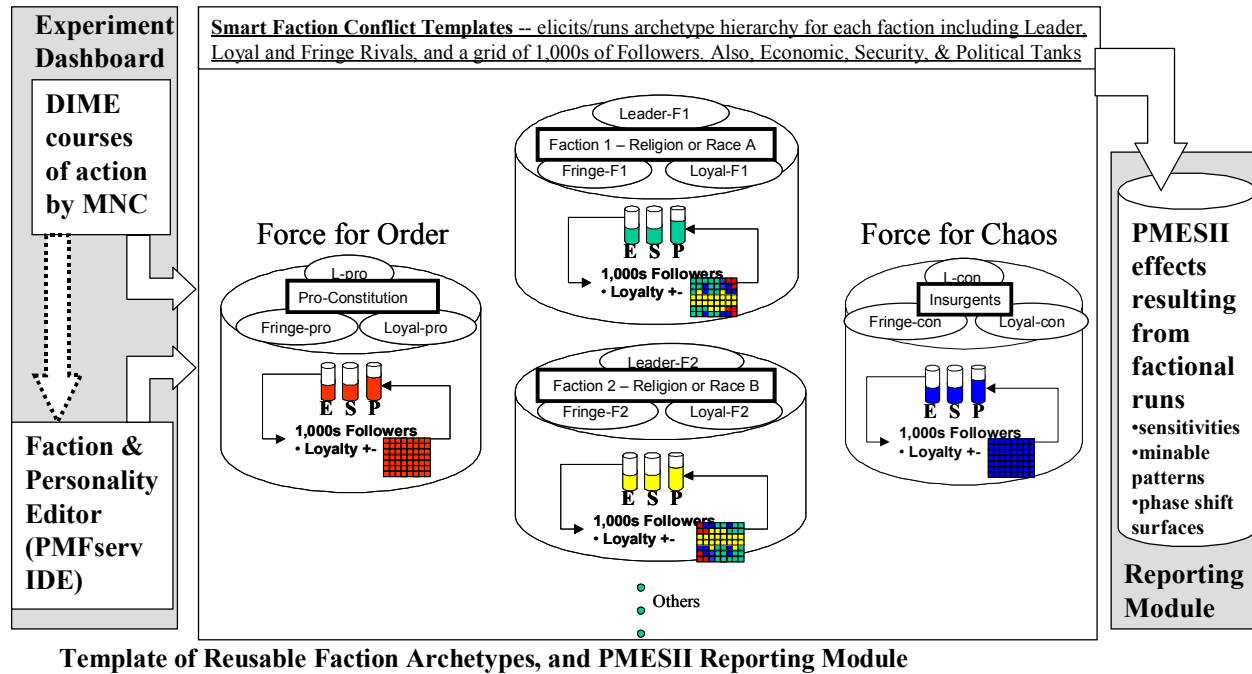
- 1) find policies that prevent conflict situations from escalating to civil strife or war: e.g., see Collier & Hoeffler (2004)
- 2) determine how best to defeat foes adept at using local PMESII effects to their own advantage, and find alternative ways to end conflict: e.g., see McCrabb & Caroli (2002).
- 3) after conflicts, manage non-kinetics operations, give aid, and help rebuild so as to avoid a return to strife: e.g., see Wood (2003); and
- 4) once we've made friends, keep friends, by avoiding our own biases (mirrors, confirmation) and exploring the factional and cultural situations that they actually face: e.g., see Heuer (2000)

If analysts and trainees are to have realistic and reliable models of the effects of DIME type operations upon PMESII dimensions, one must find ways to integrate and utilize scientific know-how across many disciplines. Part I of this article started from the bottom up – it presented a unified architecture (PMFserv) for human behavior and examined and synthesized agent-based models of leaders and followers. The focus was on what makes people join factions and commit to (or fight against) the various actions that factional leaders wish to pursue. This paper brings those components into game theory to determine whether they can help analyze such questions, and to assess the validity of the BGT approach.

2) Behavioral Game Theory (BGT) and the FactionSim Testbed

Our exploration of BGT begins by constructing a testbed (FactionSim) that facilitates the codification of alternative theories of factional interaction and the evaluation of policy alternatives. FactionSim is a tool where you set up a conflict scenario in which the factional leader and follower agents all run autonomously. You are the sole human player interacting to try and use a set of DIME actions to influence outcomes and PMESII effects. In what follows, we shall explore these issues using efforts with FactionSim testbed as illustrative examples. Section 2 begins with an examination of how FactionSim permits one to mock up general conflict scenarios, and Sections 3 and 4 then offer illustrative cases in the MidEast and in SE Asia, respectively. The reader is assumed to be familiar with Part I which delineated both the game data structures of factions that have roles to be filled (leaders, sub-faction leaders, and followers) and the available actions using different factions' resources (Economy, E, Security, S, and Politics, P) so as to influence outcomes relative to one's internal goals, standards, and preferences.

Figure 1 –Architecture for a Highly Usable FactionSim including DIME Experiment Dashboard, Smart



2.1) Definitions

Let us start by indicating that the center of Figure 3 shows there are potentially many factions, agents, and resources, however, we try to limit these to what we believe are the minimal set needed to characterize ethnopolitical factions. For an intuitive explanation, Figure 1 shows one example of a scenario template of a multi-nation or multi-faction state where there are 4 clusters of factions, each of which has a leader, two sub-factions (loyal and fringe), a set of starting resources (E-, S-, and P-Tank levels), and a representative set of over 1,000 follower agents. On the left in this example template, there is a pro-constitutional group (Pro) trying to bring order and rule of law to a region inhabited by potentially divisive factions. For scenario authoring and evaluation purposes, one can make the Pro group as strong or weak as desired by altering its tank levels relative to starting levels of each of the other factions. Likewise the form of the Pro Faction is flexibly editable as well – e.g., it can be developing or developed, repressive or democratic, corrupt or not, etc. In the example shown in Figure 1, all the institutions are still weak, so the Economic and Security tanks are not well developed. If the Pro-Faction could successfully undertake regional development there would be greater security, more employment, and more revenue to spread around – and each would be more effective. Its followers would be happier and would increase their political support, thereby raising the level of the P tank.

For rule-of-law to take sway, the Pro-Faction leader must manage his E and S tanks so as to appeal to each of the tribes or factions he wants in his alliance – the factions in the middle of Figure 1. Each of the leaders of those factions, however, will similarly manage their own E and S assets in trying to keep their sub-factions and memberships happy. After all, a high P-tank means that there are more members to recruit for security missions

and/or to train and deploy in economic ventures. So leaders often find it difficult to move to alignments and positions that are very far from the motivations of their memberships (represented by the computed level of average faction GSP in our AI model). However, even factions that hate each other are sometimes observed to form temporary alignments-of-convenience against common foes (the enemy of my enemy ...). And, of course, a leader can lose power to an opposing force such as the right side of Figure 1 which shows an insurgent force dedicated to the destruction of the Pro-Faction and hoping to replace it with an order of its own. Again, one can edit this force to whatever strength and maturity of insurgency is relevant for the scenario of interest.

These objects and attributes were more fully explained in Part I and are now listed here in general form as:

Faction = { Properties {name, identity repertoire, demographics, salience-entry, salience-exit, other}
 Alignments {alignment-matrix, relationship valence and strength, dynamic alliances}
 Roles {leader, sub-leader, loyal-follower, fringe-follower, population-member},
 Resources(R) = Set of all resources: {econ-tank, security-tank, political support-tank} }

Resource-i = { r = a resource
 $r_{r,f}$ = {Resource level for resource r owned by faction f, $r_{r,f}$ ranges from 1 to 100}
 $\Delta_{r(a,b)}$ = {Change in resource r on group a by group b} = Δ_r
 T = Time horizon for storing previous tank values
 Dev-Level = {Maturity of a resource where
 1=corrupt/dysfunctional, 3=neutral, 5= capable/effective} }

Actions(A) = { Leader-actions(target) = {Speak(seek-blessing, seek-merge, mediate, brag, threaten),
 Act(attack-security, attack-economy, invest-own-faction,
 invest-ally-faction, defend-economy, defend-security))
 Follower-actions(target) = {Go on Attacks for, Support (econ), Vote for, Join Faction,
 Agree with, Remain-Neutral, Disagree with, Vote against, Join Opposition Faction, Oppose with
 Non-Violence(Voice), Rebel-against/Fight for Opposition, Exit Faction } }

In FactionSim, these objects and attributes are all modeled with an open-architecture and are available through XML remote procedure calls. This allows 3rd party software to be plugged in or federated: (1) to mine and instantiate values from outside data sources; (2) to run someone else's more detailed models of the resources and institutions that manage them; and/or (3) to translate our agents' actions into simulators with finer levels of spatial, temporal, and/or visual detail. The only constraint is that the attributes must be transformed into and out of the representational form and units that we adopt here.

Returning to the game, there are assumed to be a set of multiple games, $G = \{G1, G2, \dots, Gn\}$, proceeding simultaneously, where each game may in fact be evolving dynamically into another form of game. For example, within a faction there might be games between rival leaders, between leaders and followers, and follower on follower. The across-faction games include attempts to cooperate and/or compete with other factions' leaders and followers, and/or attempts to defeat factions aimed at your own downfall. More precisely, the agents that populate and play FactionSim participate in a multi-stage, hierarchical, n-player game in which each class or type of agent (Dx) or simply X, observes and interacts with some limited subset Y agents (human or artificial) via one or more

communication modes. We make three (empirically plausible) assumptions about multiple hierarchies of agents, namely, that agents (1) play multiple distinct games, (2) are cognitively detailed (the agents described in Part I have approximately 600 behavioral parameters each), (3) and are self-serving and attempt to maximize their utility (u) within each iteration of a game defined as:

$$GAME = (a \in A, U_x, D_x) \text{ for } x \in X \quad [1]$$

Despite efforts at simplicity, stochastic simulation models for domains such as this rapidly become complex. If each leader has 9 action choices “on each of the other (three) leaders”, then he has $729 (= 9^3)$ action choices. Each other leader has the same, so there are 729^3 (~ 387 million) joint action choices by others. Hence the strategy space for a leader consists of all assignments of his 729 action responses to each of the 729^3 joint action choices by the other three. This yields a total strategy set with cardinality 387 million raised to 729, a number impossibly large to explore.

The DIME Experiment Dashboard consists of the boxes along the left edges that hold editors and viewers. Starting on the top left, this is the batch inputs from the player (USA or multi-national coalition) consisting of DIME courses of action (operations) that s/he thinks may influence outcomes favorably. This input can be one course of action, or a set of parameter experiments the player is curious about. On the bottom left is the editor of the personalities for the leaders and sub-leaders, and of the key parameters that define the starting conditions of each of the factions and sub-factions. Certain DIME actions by the player that are thought to alter the starting attitudes or behavior of the factions can flow between these two components – e.g., a discussion beforehand that might alter the attitudes of certain key leaders (Note: this DIME action is often attempted in settings with real SMEs and diplomats playing our various games). To summarize, here are the parameters that are available via the dashboard:

1. Tank Levels (Econ, Sec, Pol) for current turn and history including "Blame/Credit" for changes to levels.
2. Relationship Dictionary - Current relationship values from one group onto another, individuals as well.
3. GSP Weights - PMFServ GSP weights for all agents in scenario
4. State Parameters - PMFServ state values for all agents in scenario (demographic, socio-economic, etc)
5. External Action - The current external DIME courses of action(s) which analysts choose.

Agent actions are also viewable on the dashboard, but cannot be directly altered because the action choices of agents are by definition endogenous to the game(s) being played. To force them, we would make that agent external.

On the far right of Figure 1 is a module for capturing, observing, and analyzing the PMESII effects of the DIME actions. The idea is to include features to help the user visualize and understand not only the robustness of alternative policies (for avoiding conflict and enlisting cooperation), but also help to clarify the rationales of the simulated archetype participants. For example: What is driving the leader and follower decisions and choices? What is happening to their grievances, emotions, and out-group feelings?

2.2) Game Analysis

This section presents the basic game analyzed in this paper – iterated prisoner’s dilemma for multiple players. This game formulation is the simplest game one can analyze involving conflicts between (and within) factions. While it greatly over-simplifies real world conflicts as well as what is simulated in subsequent sections, it

helps to clarify many of the key elements of these conflicts. Thus this simple model serves as a building block both for BGT and for the understanding of BGT. The treatment in this section will examine how games are treated by two types of agents, namely:

- **Rational Actors:** Presumed normative as in early economic theory and intro game theory classes - perfectly informed, purely logical, and motivated by self-interest to maximize their material payoffs. All actors have identical payoff functions.
- **Descriptive agents:** Following the new tradition of BGT, these agents are characterized by descriptive rather than normative models. Also, as in Part I of this paper, we profile real world individuals with best-of-breed social science instruments. Aside from material payoffs, these agents attend to emotional, moralistic, and injustice issues; they may commit errors and use biased heuristics; and they may see games through a different lens (e.g., fast track to next life).

Starting with the simplest formulation, let us consider a game for dyadic interaction between Leaders X and Y as shown in Table 1. If X and Y were to cooperate, this might not serve their short-term goals (for example, they may lose some support from extremists within their membership). But a conflict-free world would enable these leaders to focus on socio-economic growth. For the sake of simplicity, we make no distinction between cooperate, being passive, and allying in a dyadic interaction, and hence assume that the choice of “cooperate” by both leaders will result in an alliance where the leaders share resources and obtain respective payoffs of R_{dx} and R_{dy} . Here, R_d ($= R_{dx} + R_{dy}$) represents the value of the contested resource share available.

Let us say that leaders X, Y (and Z as we will see later) have an existing relationship or level of attraction to each other and to other leaders such as Z. In any interactions between leaders, we assume that their existing dyadic relationships are altered by an amount ΔK , which is a function of relationships between the leaders as well as the actions taken. Descriptive agents’ relationships are updated by adjusting the payoff amount ΔK_{xy} to reconcile relationships that were unbalanced. For example, the extra cost for X of ruling territories outside an alliance, compared with the cost of ruling the territories in alliance with Y, could be described as the cost of loss of relationship $|\Delta K_{xy}|$. Similarly, let emV be the emotional (non-material) utility associated with taking an action. The descriptive agent receives a positive payoff, if the actions considered are in alignment with his or her value system (GSP Tree).

If a leader (X) chooses “Fight” (F) and the other leader (Y) chooses “Cooperate” (C), this will be taken to mean that X attacks Y, and that Y does not fight back (but may take shelter behind some existing defense in place). If X attacks Y with level of effort, Q_j , $j = 1, 2$, where the levels of the attack have been normalized: $0 \leq Q_j \leq 1$, then a proportion of the contested resource is transferred to the winner. If both leaders choose F then the probability of success in the ensuing battle is taken to be proportional to the relative strengths of the leaders and the efforts they put in, measured as level of attack $[P_x = Q_j \cdot R_x / R]$. The quantity lost in the battle, which is the consequence of the attack, is proportional to the target resource being contested. Therefore, expected loss in a given battle for a target is proportional to (level of attack)(relative strength of attacker)(contested resource of attacked) $= Q_j y_x (R_y / (R_x + R_y))$

Rdx. Note also that before the game is played, we assume that the leader's decisions are unknown to each other, but that each is certain about the strength with which they are attacked.

Let $CstB \geq 0$ represents the cost of staging a battle in a dyadic interaction. This is the fee that any attacker or fighter must pay. In order for this game to be identified as a prisoners' dilemma, we must have Expected Gain in battle $> |\Delta K_{xy}| + CstB$. $|\Delta K_{xy}| \geq 0$ by definition. Note that ΔK_{xy} will be positive in a compromising or allying relationship, but negative in a conflict. However, in order to emphasize the direction of costs, we are using the absolute value of the ΔK to treat it as a cost.

Consider a single shot game for a dyadic interaction. The joint choice possibilities are F/F, F/C, C/F and C/C with outcome payoffs as summarized below.

$$S2x[FxFy] = Rdx - Qjyx. Ry/R2. (Rdx) + Qjxy. Rx/R2. (Rdy) - CstB - |\Delta K_{xy}| + emV(Fx, Fy)$$

$$S2x[FxCy] = Rdx + Qjxy. Rx/R2. (Rdy) - CstB - |\Delta K_{xy}| + emV(Fx, Cy)$$

$$S2x[CxFy] = Rdx - Qjxy. Ry/R2. (Rdx) - |\Delta K_{xy}| + emV(Cx Fy)$$

$$S2x[CxCy] = Rdx + |\Delta K_{xy}| + emV(Cx, Cy)$$

$S2x[FxFy]$ is the payoff for X in a dyadic interaction (indicated by Scenario 2 or S2), when both X and Y are fighting. Similarly, $S2x[FxCy]$ is the payoff for X when X is fighting and Y is compromising or cooperating (read as being passive and not fighting back). Note that only the payoffs for X are given, but in a dyadic interaction, one can obtain the payoff for Y by symmetry. i.e. $S2x[FxCy]$ and $S2y[CxFy]$ are identical in structure. Numerical values may differ due to individual differences. With single shot as well as finitely repeated games, one subgame perfect equilibrium is mutual fighting (Dutta, 2000). Rawls deficient (Macy, 2006) mutual conflict, however, is too myopic for repeated games. With infinite horizons several alternatives (based on different subgame perfect equilibria) emerge. For example, mutual cooperation can constitute an equilibrium in an infinitely repeated game (e.g., see Folk Theorem in Dutta, Kaneko (1982)), provided it is Pareto optimum and a rule such as tit-for-tat or grim-trigger is established to counter the temptation of unilateral defection. For mutual compromise to be a Pareto Optimum, we should have:

$$S2.4x[CxCy](1+i)/i > S2.2x[FxCy] + S2.1x[FxFy](1/i) \text{ and } S2.4y[CxCy](1+i)/i > S2.3y[CxFy] + S2.1y[FxFy](1/i)$$

It follows that by symmetry:

$$S2.4x[CxCy](1+i)/i > S2.3x[CxFy] + S2.1x[FxFy](1/i) \text{ as } S2.2x[FxCy] > S2.3x[CxFy]$$

$$S2.4y[CxCy](1+i)/i > S2.3y[CxFy] + S2.1y[FxFy](1/i) \text{ as } S2.2y[FxCy] > S2.3y[CxFy]$$

Table 1 – Two Leader (Dyadic) Outcomes for the Iterated Prisoners Dilemma Game

<u>Equilibrium Conditions:</u> There are many equilibria in the infinitely repeated game. Pareto optimum mutual cooperation is also rational as the horizon increases and the mechanism to punish defection is effective. For this case we need: $S2.4x[CxCy](1+i)/i > S2.2x[FxCy]+ S2.1x[FxFy](1/i)$ and	<u>Payoffs for repeated game with infinite horizons:</u> <table><tr><td rowspan="2"></td><td colspan="2">Leader Y</td></tr><tr><td>Fight</td><td>Cooperate</td></tr></table>		Leader Y		Fight	Cooperate
	Leader Y					
	Fight	Cooperate				

$S2.4y[CxCy](1+i)/i > S2.3y[CxFy] + S2.1y[FxFy](1/i)$ It follows by symmetry that: $S2.4x[CxCy](1+i)/i > S2.3x[CxFy] + S2.1x[FxFy](1/i)$ as $S2.2x[FxCy] > S2.3x[CxFy]$ $S2.4y[CxCy](1+i)/i > S2.3y[CxFy] + S2.1y[FxFy](1/i)$ as $S2.2y[FxCy] > S2.3y[CxFy]$	Leader X	Fight	$S2.1x[FxFy](1+i)/i,$ $S2.1y[FxFy](1+i)/i$	$S2.2x[FxCy] +$ $S2.1x[FxFy](1/i),$ $S2.2y[FxCy] +$ $S2.1y[FxFy](1/i)$
		Coop	$S2.3x[CxFy] +$ $S2.1x[FxFy](1/i),$ $S2.3y[CxFy] +$ $S2.1y[FxFy](1/i)$	$S2.4x[CxCy](1+i)/i,$ $S2.4y[CxCy](1+i)/i$

When examining payoffs such as these, the reader should keep in mind the distinctions between rational and descriptive agents. The best-of-breed descriptive models of agent value systems and a description of how relationships and emotional payoffs are computed were presented in detail in Part I of this article. Here we depart somewhat from that description by simplifying the differences between rational and descriptive agents to include only whether or not they attend to ΔK_{xy} and emV issues (rational actors are assumed to ignore these terms). Similarly, this formulation glosses over the issues of how actors compute the size of an attack (Q_j), how they discount (i), and how much they are willing to pay for their gambits ($CstB$) – one wouldn't even expect to use the same formulas for normative vs. descriptive computations. With that as background, we can summarize the possible game results as follows:

- **Rational Actors:** Mutual conflict or fight-fight is a well-known Nash equilibrium. We know also that if $S2x[CxCy] > \{S2x[FxCy]\}$, where x in X, Y , then mutual cooperation is Pareto optimal. However, whether mutual compromise is also a Nash equilibrium depends on the value of fighting while the other cooperates.
- **Descriptive agents:** Here the outcomes are less clear. If both agents have significant payoffs for non-violence, for example, mutual cooperation can become both a Nash equilibrium and Pareto optimum. If only one agent has such payoffs for non-violence or is too weak or poorly organized to fight, then Fight/Cooperate could become an equilibrium. If continued across iterations, this could signify genocide or ethnic cleansing. It can also account for loss of will to fight (e.g., US in Vietnam, Russians in Afganistan, Shias under Saddam). For the same reasons, mutual fighting may be an equilibrium for totally different reasons such as choosing to fight and die (the martyrdom game).

Once the single shot game described above plays out, a repeated game ensues with information from past events. For rational actors, subgame perfect equilibrium such as mutual compromise or mutual conflict, once established, tends to persist in the absence of exogenous shocks. If we assume a constant future discount rate (rate of time preference) of i for both parties, then over infinite horizons, one may write the payoffs as in Table 1.

Triadic Games (Single Shot) – In most multi-agent encounters, the outcomes will depend on the interactions between more than two agents. For example, let us consider a third player Z , who is either a single entity or in a strong dyadic alliance such that Z could be regarded as a single entity. Agent Z could interact with the

first dyad by forming an alliance or by attacking. Table 2 shows the six standard scenarios for the triadic game (rows labeled S3.1 to 3.6) and the conditions where equilibria should be expected. It is worth distinguishing how rational vs. descriptive actors would be expected to behave.

Specifically, for Rational Actors, mutual conflict or fight-fight (S3.1) is a Nash equilibrium, as cooperating while the other players attack has a low payoff. However, this may not be a Nash equilibrium if the cost of staging a battle is very high, or the agent is too poor to pay $CstB$, or attacking players have very little disputed resources at stake. In other words, the question of whether it is economically attractive to fight depends on the ability to pay for the battle from the gains. To engage in battle with both Y and Z, X must have: $Qjxy(Rx/R3) Rdy + Qjxz(Rx/R3)Rdz - 2CstB \geq 0$. To engage in battle with at least one aggressor (say Z), X must have: $Qjxz(Rx/R3)Rdz - CstB \geq 0$. For these reasons, it is highly unlikely for S3.2 to become an equilibrium where one agent passively cooperates while another attacks. Likewise, based only on material pay offs, aggressive alliances may form (S3.5), but will be unstable or exhibit weaker aggression than expected. Two leaders may choose a coalition to protect them against the strongest leader (S3.4); but this will be an unstable equilibrium, remaining in place only as long as a threat of the third leader exists. Finally, even when mutual cooperation (S3.6) is a Pareto optimum, it is not individually-rational, particularly for the strongest leader. For “peace in the world” to be a Nash equilibrium, the payoff to every player has to be greater than payoff received from any configuration of aggression, while others maintain cooperation.

By contrast, we see a different set of equilibrium possibilities for Descriptive Agents. For example, “Fight-Fight” (S3.1) may not be an equilibrium for the reasons discussed under the dyadic game. Remaining passive while the opponents are aggressive could be preferred (S3.2, S3.3), if the expected cost of staging a battle and its non-material consequences are higher than remaining passive and taking a hit. For example, an agent might be strongly conditioned to be “non-violent” due to its value system, or is too risk averse to pay a high cost of staging a battle (prospect theory). Therefore, with Descriptive Agents, the passive nature of the agents and their standards may alter the fact that mutual conflict is a Nash equilibrium. Likewise, positive relationships and emotional payoffs will encourage the formation of coalitions such as in S3.4 or S3.5, just as adverse relationships and grievances may prevent the formation of other coalitions. There can also be agent X with a grievance that may engage in fighting even when $Qjxz(Rx/R3) (Rdz) - CstB < 0$. Likewise, one may find that coalitions are made up of enemies who return to fighting each other once a common threat is overcome (S3.5). This means that mutual cooperation (S3.6) is equally difficult to predict, since when the grievances are significant (e.g. $emV(Cx, Fz)$ is high), the emotional cost of cooperating is high, and there is an emotional incentive to fight

Table 2 – Nash Equilibria Conditions for the Scenarios of a Three Leader (One Shot) Game

Scenario	Conditions for Equilibrium
S3.1: [FxFy, FxFz, yFz] All Fighting	When all three are fighting each other, the payoff for any leader X may be expressed as: $S3.1x = Rdx - Qjyx(Ry/R3) Rdx + Qjxy(Rx/R3)Rdy - Qjzx(Rz/R3) Rdx + Qjxz(Rx/R3) (Rdz) - 2CstB - \Delta Kxy - \Delta Kxz + emV(Fx, Fz) + emV(Fx, Fy).$
S3.2: [Cx Fy, Cx Fz, Cy Cz]	Payoff for X for remaining passive: $S3.2x = Rdx - Qjyx(Ry/R3)Rdx - Qjzx(Rz/R3)Rdx + \Delta Kxy + \Delta Kxz + emV(Cx, Fz)$

aggressors y and z independently attack a passive x	$+ emV(Cx, Fy)$ Payoff for any one aggressor, assuming aggressors Z and Y do not interact: $S3.2z = Rdz + Qjzx (Rz/R3)Rdx/2 - \Delta Kxy - CstB + emV(Fz, Cx)$
S3.3: [CxCy, CxFz, CyFz] Sole leader attacks as peaceful coalition who cooperates	The payoff for any of the targets X, Y: $S3.3x = Rdx - Qjzx(Rz/R3)(Rdx + Rdy)/2 - \Delta Kxz + \Delta Kxy + emV(Cx, Fz) + emV(Cx, Cy)$ Payoff for the aggressor: $S3.3z = Rcz + Qjz_xy(Rz/R3)(Rdx + Rdy)/2 - 2CstB - \Delta Kzx - \Delta Kzy + emV(Fz, Cx) + emV(Fz, Cy)$ Even with structural advantages brought about by coalition formation (as opposed to S3.2), if the coalition still remains passive under attack, this scenario is even more accentuated. Therefore, this scenario is even more unlikely to be a Nash equilibrium for Rational Actors except as mentioned in S3.1.
S3.4 & S3.5: [CxCy, FxFz, FyFz] Individual leader Vs coalition fighting	Scenarios S3.4 and S3.5, respectively, deal with the situations where a third aggressor attacks a coalition who does fight back, and aggressors in alliance (compromise/ coalition) attack a target who also fights back. Since there is no first mover advantage, by symmetry, these scenarios are identical and we only present S3.4. The payoff for any of the targets X, Y who fight back as a coalition are: $S3.4x = Rdx - Qjzx(Rz/R3)(Rdx + Rdy)/2 + ((Qjxz Rx + Qjyz Ry)/R3)(Rcz)/2 - CstB/2 - \Delta Kxz + \Delta Kxy + emV(Cx, Fz) + emV(Cx, Cy)$ If a leader in a coalition does not fight back an attacker, his or her utilities are $Rdx - Qjzx (Rz/R3) (Rdx) - \Delta Kxz + \Delta Kxy + emV(Cx Fz) + emV(Cx Cy).$ The payoff for the aggressor is: $S3.4z = Rcz + Qjz_xy(Rz/R3)(Rdx + Rdy)/2 - 2CstB - \Delta Kzx - \Delta Kzy + emV(Fz, Fx) + emV(Fz, Fy)$ An advantage of defending an attack through an alliance that is already in place is obvious, as it involves higher likelihood of success, cost savings in terms of the costs of staging battle, and stronger relationships. However, an alliance also reduces the spoils of the fight, as it gets shared between the members of the coalition. Therefore, as long as CstB/2 does not out-weigh the benefits, there is incentive to form coalitions. Although significant disparity in assets may reduce the benefits of the coalition for the stronger leader, the effort invested in a battle also decided how much is contributed to the common cause of aggression. A leader may also attempt to free-ride or reduce the load by lowering the effort put into the battle. Also, there is a disincentive to attack a coalition as opposed to individual leaders, unless the attacking leader has a resource advantage.
S3.6: [CxCy, Cx Cz, Cy Cz]	The payoff from mutual cooperation to any player X is:

Mutual Cooperation	$S3.6x = Rdx + \Delta K_{xy} + \Delta K_{xz} + emV(Cx, Zc) + emV(Cx, Cy).$ Mutual cooperation (S3.6) is Pareto optimal with payoffs larger than the Nash equilibrium, however, if there is significant disparity between the resources (Rx) and the disputed resources (Rdx), there is a strong incentive for a stronger leader to start an aggressive move. Under these circumstances, mutual compromise is not a Pareto optimum. For the strongest leader, say Z, not to start aggression unilaterally against x, we have: $Rdz + \Delta K_{zx} + \Delta K_{zy} + emV(Cz, Cx) + emV(Cz, Cy) > Rdx + Qjzx (Rz/R2) Rdx - \Delta K_{zx} + \Delta K_{zy} - CstB + emV(Fz, Cx) + emV(Cz, Cy).$ Similarly, a coalition of two (x and y) to refrain from attacking the leader Z, who is currently in the same coalition: $Rdx + \Delta K_{xz} + \Delta K_{xy} + emV(Cx, Zc) + emV(Cx, Cy) > Rdx - Qjzx(Rz/R3)(Rdx + Rdy)/2 + (Qjxz Rx + Qjyz. Ry)/R3)((Rcz)/2 - CstB/2 - \Delta K_{xz} + \Delta K_{xy} + emV(Cx, Fz) + emV(Cx, Cy)$
--------------------	---

Repeated Triadic Games – Table 2's equilibria may shift if the game is repeated. For example,

- Mutual cooperation, in the complete absence of memory, is unstable, in part since maintaining relationships involves memory. For Rational Actors without memory, one may assume S3.1 is the single shot (subgame perfect) equilibria with positive payoff cycles, and it will be repeated.
- If the resource levels of the leaders are comparable, mutual fighting will be a rational and subgame perfect equilibrium in the repeated game (even when mutual compromise is the Pareto Optimum).
- Coalition of two leaders against the strongest leader is a subgame perfect equilibrium, when there is disparity in resources, and the resources of the second and the third strongest leaders are comparable. Under these conditions, a bipolar world could evolve, albeit temporarily. This is in congruence with Landscape Theory of Aggregation (LTA) (Axelrod and Bennett, 1993), which demonstrates that a bipolar (two factions) configuration is stable (Nash equilibrium) for a collection of antagonistic states.

Repeated games might lead to relatively predictable phases of agent learning and new equilibria emerging in a multi-stage fashion. For example, if a triad starts with disparity of resources so Z is the strongest while X and Y have comparable resources. The payoffs for the infinite horizon game can be derived for triadic interaction in a fashion similar to dyadic interaction shown earlier by dividing by $(1+i)$. The following sequence of behavioral cycles is an example selected for its interesting combination of a variety of subgames. Here, we assume that a coalition of two leaders (S3.4x) act against the strongest leader (S3.4z), which will be repeated until the relative powers change due to repeated conflicts. The powers equalize in T1 turns, and the mutual fighting continues till time T2, when the resources available are too little to pay for further battles. At this stage, mutual cooperation is the only subgame perfect equilibrium, but it is unlikely to remain stable:

$$\text{PAYOFF}_x = \sum_{t=0}^{T_1} \frac{S_{3.4}x(t)}{(1+i)^t} + \sum_{t=T_1+1}^{T_2} \frac{S_{3.1}x(t)}{(1+i)^t} + \sum_{t=T_2}^{T_3} \frac{S_{3.6}x(t)}{(1+i)^t} \quad [2]$$

Two leaders x and y against the strongest z	All leaders with comparable strength, mutually fighting.	Exhausted leaders exhibiting mutual cooperation
---	--	---

Where $S_{_}x(t)$ refers to the payoff for x in scenario $S_{_}$ occurring in time step t.

Rational Actors with Memory – In repeated games, if the agent histories are remembered, no agent is excessively powerful, and agents start with mutual cooperation, then the following is the well-known mixed strategy that will prevail: attack if provoked (tit-for-tat) to deter other leaders from taking advantage, but otherwise cooperate. Thus long periods of cooperation punctuated by occasional conflicts may occur. Ignoring rare periods of conflict, one may write the payoffs for any given agent as:

$$\text{PAYOFF}_x = \sum_{t=0}^T \frac{S_{3.6}x(t)}{(1+i)^t} \quad \text{Over infinite horizons } (T \rightarrow \text{infinity}), \text{ PAYOFF}_x = S_{3.6}x(t)/i \quad [3]$$

Descriptive Agents – Here we can make no clear predictions, barring insight into the nature of the individual agents' preferences, standards, grievances, and personalities. Mutual cooperation would be a Nash equilibrium if strong "non-violent" or pacifist values (emotional factors) and/ or positive relationships are sufficient to keep the parties from slipping into conflict (to exploit short-term advantages). If those leaders' total and the disputed resources are comparable, cooperation would persist in the absence of exogenous shocks. Likewise, strong grievances between the leaders, lack of trust, and negative relationships (all of which capture history/memory) could also push the parties into conflict, even when they're not strong enough to gain success in battles. Also, when relative peace exists, if there are no values or long term relationships to hold them back, descriptive players can easily be swayed by occasional transgressions. Such aggressions will be remembered and agents will want to "settle grievance scores", resulting in a spiraling of conflicts. While material benefits of mutual cooperation under relative balance of power with trust are often sufficient for rational players to cooperate, the task of motivating Descriptive Agents to cooperate will be more difficult, as we will see later in the simulations.

In summary, this section has highlighted some of the key factors that influence outcomes in dyadic and triadic games. While this could of course be expanded to four and more leaders, the above discussion suffices for our present purposes. Hence we turn now to the simulation of two real-world factional conflicts and to the analysis of these simulation outputs. Since many games are going on simultaneously in FactionSim, we must be cognizant of an array of possibilities. Section 3 examines how alternative DIME actions might alter game outcomes for a two dyad, 4-leader game. Section 4 then extends this to a leader-follower game.

3) FactionSim: Studies for the MidEast Today

During the spring semester of 2006, students in a graduate course taught by the lead author took the FactionSim (and PMFserv) artifacts and engineered a number of reusable factions for simulations involving the MidEast. The course was taught under the 'Coop-Coop' pedagogy in which the students were organized into teams.

There were 5 four-person teams, where the task of each team was to build at least one faction model. Each team member was responsible for one of four specialties within their team: data collectors/ analyzers, PMFserv knowledge engineers, programmers, and integrators. These specialists in turn also participated in specialty teams for learning and for sharing resources and knowledge sources. Each of the authors of this paper coached a specialty team. The point of this teaching model is that students learn by discovery and by teaching each other. They also learn skills they can readily transfer to the workplace since they must cooperate across specialty teams, with members of their faction team, and across faction teams to assemble a larger whole than any one team alone could create. Evaluations at the end of the semester indicated this was a popular teaching model.

The five student teams assembled a total of 21 PMFserv leader profiles across 7 real world factions so that each faction had a leader and two sub-faction leaders. The seven factions – government (2 versions - CentralGov and LocalGov), Shia (2 tribes), Sunnis, Kurds, and Insurgents – could be deployed in different combinations for different scenarios or vignettes.

The leader and group profiles were assembled from strictly open source material and followed a rigorous methodology for collecting evidence, weighing evidence, considering competing and incomplete evidence [see for example Bharathy (2006)]. Popular sources across all groups included Brookings' Iraq Index, CIA Factbook, news archives such as BBC, CNN, Fox, and Al Jazeera, and less trusted sources such as Wikipedia and a number of Iraqi blog sites. Conveniently, the Brookings Iraqi Index is organized exactly according to our three resource tanks: economy, security, and politics. Each team also uncovered in-depth studies and reports from think tanks as well as books and journal articles that helped to inform their models [see Sageman (2005), among others]. All of this was translated into one of three types of data/ empirical information employed in their models:

- Numerical data as well as empirical materials on Iraqi factions, particularly the violent incidents,
- Anecdotal information and quotes from interviews about the decisions made, along with the contexts of these decisions, by the specific personnel being modeled, and
- Culture specific information for the factions from such studies as GLOBE (House et.al., 2005), as well as religious doctrines affecting the people of concern, and published political platforms of the diverse sub-factions.

Each faction team produced technical reports detailing data collected into evidence tables and describing (i) how evidence was handled, (ii) alternative interpretations of that evidence for calibrating parameters, (iii) documentation of the models they produced and (iv) rationale behind various other parameters and markups. Also, each group produced and tested their own factional leader models, including parameter-tuning results to ensure that the leaders behave on test cases as do their counterparts in the real world. Rather than a validation against test data set, this was designed more to tune against the training data. Two of these leaders' GSP trees were shown in Figure 2 of Part I to illustrate how the primary opponents (government vs. insurgency) differ in terms of attributes like ingroup bias, aggressive attitudes toward outgroups, willingness to conduct asymmetric attacks, humanitarianism, range of scope, and power need, etc.

3.1) Validity Assessment

The primary validity assessment test of these models relied upon a DARPA-sponsored effort labeled “Integrated Battle Command” which paid for 3 teams of subject matter experts to play the Multi-National Coalition (MNC) and evaluate the models and outcomes of DIME actions taken for various vignettes. These sessions were run at Joint Forces Command for 2 weeks in May 2006 after the end of the semester. The specific leaders and factions (together with any tuning of these agents) included vignette inputs, MNC courses of action attempted, and model outputs that were subsequently designated as classified material. Hence these results cannot be presented here, though the idealized runs presented in the next section should be sufficient to give readers a sense of what transpired. The SMEs ranged across areas of military, diplomatic, intel, and other PMESII systems expertise. Within each vignette the SMEs attempted dozens of courses of action across the spectrum of DIME possibilities (rewards, threats, etc.). One interesting COA is reflected in earlier Figure 1 by the vertical arrow on the left of the chart linking the MNC to the Personality Editor. That is, a popular COA of the diplomats was to ‘sit down’ with some of the persuadable leaders and have a strong talk with them. This was simulated by the senior diplomat adjusting that leader’s personality weights (e.g., scope of doing good, treatment of outgroups, etc.) to be what he thought might occur after a call from President Bush or some other influential leader.

The SME team playing the MNC presented their opinions at the end of each vignette. The feedback indicated that the leader and factional models corresponded with SME knowledge of their real-life counterparts, and (predictably) specific recommendations were offered for improving the realism and detail of the resources and institutions modeled by our simple ESP tanks. No further comments will be offered here about the Turing test used. For purposes of illustration, and to facilitate discussion about analysis of outcomes, the simulation results in the next section make use only of the open-source models created by the students with simple DIME courses of action.

3.2) Experiment #1: Elasticity of Conflict in Iraq Due to Outside Support

This section shows runs of 4 factions initially organized into two weak alliances (dyads): (i) CentralGov trying to be secular and democratic with a Shia tribe squarely in their alliance but also trying to embrace all tribes, (ii) a Shia tribe that initially starts in the CentralGov’s dyad but has fundamentalist tendencies, (iii) a secular Sunni tribe that mildly resents CentralGov but does not include revengists, and (iv) Insurgents with an Arab leader trying to attract Sunnis and block Shia control. Each faction has a leader with two rival sub-leaders (loyal and fringe) and followers as in Figure 1 – all 12 are named individuals, many are known in the US. This is a setup that should mimic some of the factional behaviors going on in Iraq, although there are dozens of political factions there in actuality.

Figure 2 reveals the outcomes of three sample runs with these factions. Since this paper focuses on conflict dynamics (prevention, termination, etc.), we omit discussion of dynamics affecting the economy and political tanks to save space. The left hand column shows activity over time (a tick of the simulator is one week) that happens in the Security Tanks. The vertical axis indicates the normalized fraction of the sum across all security tanks in these factions, and thus the strip chart indicates the portion of the sum that belongs to each faction. Rises and dips correspond either to recruiting and/or battle outcomes between groups. One can inspect the right hand column to get a sense of which factions carried out positive, neutral, or negative acts toward other factions - this view has no time dimension. Negative actions (histograms above a minus sign) show attacks on another group of a given color.

Figure 2a is a run where different outsiders take DIME actions that support the CentralGov and the Insurgents. The Economic and Security Tanks of CentralGov and the Security Tank of Insurgents receive a 1% boost every tick of this run (slight advantage to CentralGov). Here we can see that the CentralGov is unable to establish military dominance, and conflict is fairly continuous among the groups. The Insurgents are weakened by military attacks from the Shia and their own failed attacks on the CentralGov, which concentrates on defense. CentralGov allies with the Shia, providing Economic Aid and sponsoring them to conduct more of the attacks. The Sunni focus on opportunistically culling support from whomever seems most powerful at the moment (notice token + actions on right histogram), but also occasionally retaliate against each of the other factions following collateral damage to the Sunnis. A take-away lesson of this run seems to be that democracy needs major and continuous outside help, as well as luck in battle outcomes and some goodwill from tribes for it to take root.

In Figure 2b, outside support for CentralGov is removed (the coalition pulls out), and one can see they collapse precipitously. The Insurgents attack CentralGov, maintaining strong military pressure to keep CentralGov weak. Meanwhile, the Shia chip away at the Insurgents. Shia support and recruitment grow, leading them to battle the Insurgents to what looks like a steady back and forth for the second year of the run. Sunni-secular is not entirely innocent throughout this, opportunistically attacking the weakened CentralGov while also spreading tokens of support to gain favor from the Shia alliance. The Insurgents are too distracted by the outside alliance to worry about counteracting any such flirtations within their own alliance. The take-away lesson seems to be that civil war rules.

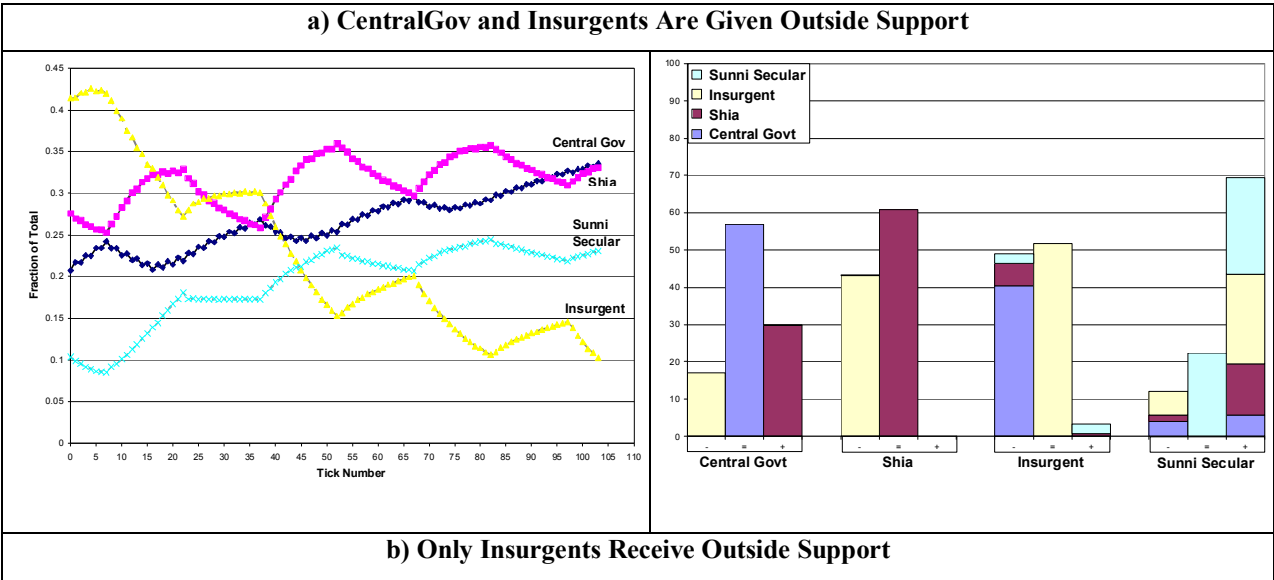
Finally, in Figure 2c, we also remove support for the Insurgents, a DIME action equivalent to shutting the borders so no outside recruitment or funding is possible. The Insurgency is fairly rapidly dispensed with, but once again the outcome seems to reflect tribal division of spoils, rather than emergence of a democratic nation. The CentralGov is permitted to remain though they are constantly paying off the Shia (right side histograms) and likely are a fundamentalist puppet. Once on top, the Shia also gain token support from the Sunni, who have built themselves up while the others fought to the point where they are strong and pretty much left alone. And, to be sure, they always tend to provide token support to the stronger side. A take-away lesson from this run seems to be that factions in this part of the world, when left to their own devices, will resort to tribal division of spoils.

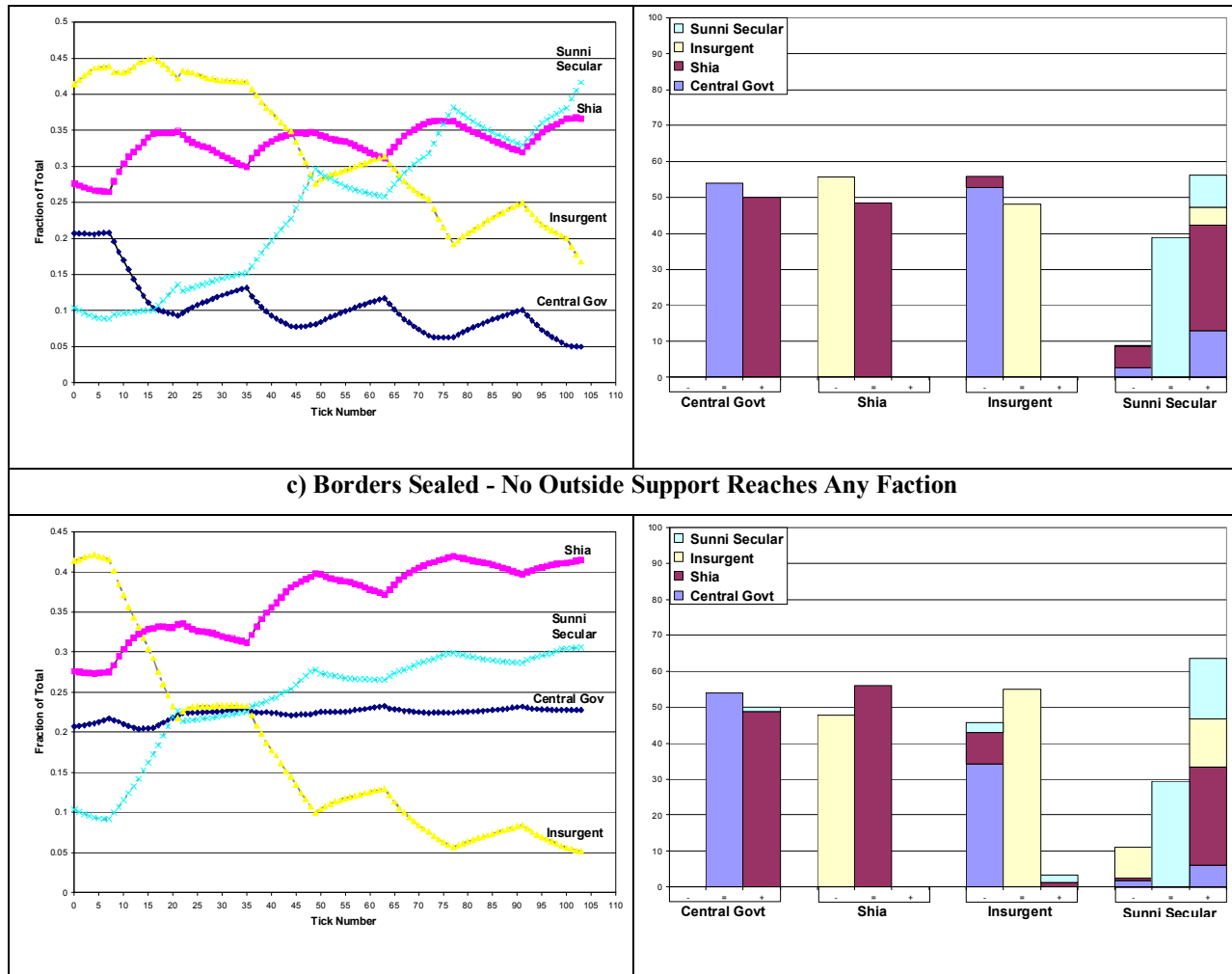
To sum up, these three sets of runs serve as a nice example of the quandary for those trying to formally summarize the outcomes from human behavior model-infused EBO simulators. On one level we might be tempted to simplify the summary of outcomes across cases by agreeing with the intuitive “take-aways” of the three runs, respectively – democracy is hard to achieve, civil war is likely, and tribalism is prevalent. No surprises in any of those. However, a secondary goal of this research was to explore whether game theory, as formulated in Section 2.2, could help to guide DIME policy choices and make sense of the outcomes or, if not, what might be some research goals for improving our game theory framework?

In the game theory formulation of Section 2.2, we made several terms explicit that can now be seen to be vital extensions to a standard payoff formulation. Let us see how just two of these work -- relationship changes (ΔK) and emotional factors (emV) -- two PMFserv factor sets that have no pre-scripted rules governing their dynamics. Though assumed away by textbook game theory or Nash Equilibria, as would be expected, factions in our runs are more likely to attack groups with which they have negative relationships and strong emotions. Relationship and

emotions also factor into the formation of alliances. For example, across all runs, CentralGov has a friendly relationship towards the Shia, who are moderately positive back. This leads to CentralGov giving aid to the Shia and consistently forming an ally. Likewise the Sunni Secular have slight positive feelings towards the Insurgents, and are more likely to assist them, unless others are more powerful. In an asymmetric world, actions have spillover effects across factional relationships as well. In particular, attacks will have collateral damage that affects groups of the same sect (a surrogate for the reality of having to attack opponents fighting from within civilian districts). This affects the Sunni Secular, who receives collateral damage from Shia attacks on the Insurgents. As a result, the Sunni Secular sometimes attack the Shia group. Alliance also can result in extra relationship damage. If an ally supports a group under attack and loses troops, it will sour towards the instigator. The Shia alliance with CentralGovt causes them additional hatred towards the Insurgents for this reason. Finally, some action choices seem to have purely emotional payoffs. For example, from an economic perspective, the payoff from attacking an enemy with zero economy is zero - a wasted turn. Yet in these runs, when the Insurgents fail, the Shia still occasionally attack them simply because the Insurgents are their enemy. Emotional payoffs are at least as important as material ones, but the two scales are inherently difficult to compare. Thus the ΔK and emV terms of Section 2.2's payoffs are vital to game theoretic formulations that a human behavior model such as PMFserv is able to help the analyst to generate and understand. We omit the strip charts of changing ΔK and emV strengths due to page limits.

Figure 2 – Conflict and Cooperation in Iraqi Factions Under Alternate DIME actions (mean of 100 runs).





Another question is whether Nash Equilibria (or Pareto Optima) predicted from Table 2 emerged and how can this help to guide DIME policy choices? Frequently, a leader will take an action multiple times in sequence, indicating a temporary equilibrium for that leader's decision process. Unfortunately, these stable periods are generally short in length. This problem increases as a function of the number of leaders in the scenario, as a single leader can disrupt an equilibrium. Rational actor equilibria (equation [3]) only appears in the Sunni tribe, who retaliate when attack, but who primarily cull favor and cooperate, leading to their steady improvement in most scenarios. Of course this is not the revengist Sunnis, nor the Iraqi Mafia, but a relatively peaceable clan. All other groups are conflict-ridden and only run 3 (Figure 2c) – the run where minimalist DIME interference occurs – seems to lead to the emergence of a stable equilibrium, albeit a puppet fundamentalist government.

4) Impact of Leader Action on Follower Choices: FactionSim for SE Asia

The previous section explored conflict vs. cooperation between the leaders of different factions but ignored whether the followers were going along or resisting those actions. In this Section, by contrast, we examine the decisions of followers to cooperate with or fight against their factional leaders. Without naming the actual country or

1
2
3
4 leaders and in keeping with our game notation, we shall refer to the Bhuddist majority as X and their leader as LX.
5 During the 1990s, the country was relatively stable, however, in the last few years, the rural provinces have seen a
6 rise of Muslim anger against the central X government, and the internal security situation in these provinces has
7 rapidly decayed. During 2004, a small group of fundamentalist Muslims (Z) have committed an increasing number
8 of violent acts against Budhists (X) as part of a movement for a separate fundamentalist state. The level and
9 sophistication of the attacks has been increasing to the point where people are questioning whether there may be
10 outsiders assisting this group. The main policy concern here will be to answer question types 1 and 4 from the
11 introduction: how should LX address this problem so as to prevent a full blown insurgency from being spawned?
12 Why is violence rising in a region that was formerly friendly and peaceful? What are the consequences for domestic
13 politics? What would be the best targets and times to intervene?

14
15 The details and statistical evaluation of this case study have been presented fully in Silverman et al (2006)
16 and (2007), respectively, and we primarily focus here on the decisions of the followers for cooperation vs. fighting
17 so as to illuminate our discussion of FactionSim. In brief, the real world (open source) data shows, the reaction of
18 the LX to the violent incidents has been generally viewed as heavy-handed, and even inappropriate. LX has branded
19 the separatists as bandits, and has sent the worst behaving police from the north (X Land) to handle all protesters in
20 the Muslim provinces. There are many accounts of police brutality and civilian deaths and we classified the violent
21 incidents in the country based on the size and intensity of the incident. The incidents were aggregated and plotted
22 against time to obtain a longitudinal plot of incidents. The data was then longitudinally separated into ‘independent
23 sets’ with training set consisting of Jan-June 2004 while test set beginning in July 2004 and running till Dec 2004
24 ending just before the tsunami. In December 2004, the Tsunami hit and ravaged portions of these provinces. The
25 massive arrival of relief workers lead to an interruption of hostilities, but these resumed in mid-2005, and LX
26 declared martial law over the provinces in the summer of 2005.

27
28 Training data and evidence were used to calibrate three types of agents in PMFserv:

- 29
30
31
32
33
34
35
36
37
38
39 • Leader X (LX) (structure of his GSP trees are in part 1, Fig 2) - data indicates harsh, cruel, task, corrupt,
40 wealthy, successful. Sends worst behaving cops down to provinces, never discourages brutality.
41
42 • Moderate Y Followers - Lack of cultural freedom, schools, etc. Mostly rural family members who want own
43 land and autonomy.
44
45 • Radical Y Followers – tend to be sons of Moderate Y Followers who were Wahhabi and college-trained,
46 unemployed, running religious schools in family homes. Earlier Figure 3 of Part I shows the GSP trees of this
47 follower archetype.
48
49

50
51 In order to adequately populate the factional groups in FactionSim, we created X consisting of LX with the
52 just a Security tank and no other members of X’s faction (just its leader and his security tank). Next we set up
53 Faction Y with a moderate leader and the two types of followers mentioned above. Finally, Z was set up as just a
54 stimuli that periodically attacks X. The larger population of Y was run via a version of cellular automata that is
55 known as the Civil Violence model (Epstein et al., 2001), though Leader Legitimacy was replaced with PMFserv
56 agents’ view of membership. The Civil Violence model involves two categories of actors, namely villagers (or
57 simply agents) and cops. ‘Agents’ are members of the general population of Y and may be actively rebellious or not,
58
59
60
61
62
63
64
65

depending on their grievances. ‘Cops’ are the security tank forces of the Leader of X, who seek out and arrest actively rebellious agents. The main purpose of introducing the Civil Violence model is to provide a social network for the cognitively detailed PMFserv followers to interact with. The social network consists of one layer of the normal arena or neighborhoods as well as a second layer of secret meeting places, simply represented as a school. Civil Violence agents can exist in more than one layer (namely in the normal as well as school layers), however, the PMFserv agents that show up in the school layer are only the young Wahhabi- and college-trained males (Radical Y Followers). Overall, we

The bridge between PMFserv and Civil Violence includes Leader X’s orders and 160 villagers, and works as follows. LX examines the state of the world and makes action decisions to assist or suppress Z or Y (e.g., pay for Buddhist schools, add more cops, reduce cop brutality, etc.). The 160 PMFserv agents then assess their view of the world, react to how cops handle protester events, how their GSPs are being satisfied or not by leader actions, and to their emotional construals. The grievance level and group membership decisions by 160 archetypical villagers in PMFServ are passed via an XML bridge to 160 agents they control in the cellular automata based population model. These agents influence the neutrals of the population who spread news and form their own view of the situation. The number of Civil Violence villagers in each level of grievance (neutral through Fight Back as shown in the rows of Table 4) are added up and this information is passed back to PMFserv to help determine its starting level of grievance for the next cycle of reactions to XL’s actions. The left side of Table 4 shows the starting values as percent of population of Y that occupies each Grievance State. We discuss the right side in Section 4.1.

Table 4 –Faction Y Shifting from Relatively Cooperative (GS0-2) to Largely Fighting (GS3 & GS4)

Starting State (Avg of Weeks 1 & 2) Muslim Population at Start Is Neutral with Few Grievances Registering		End State (Avg of Weeks 103, 104) Muslim Population Reflects Radicalization and Spread of NonViolent and Violent Protest	
GrievanceState0 - Neutral	30%	6%	GrievanceState0 - Neutral
GrievanceState1 - Disagree	55%	1%	GrievanceState1 - Disagree
GrievanceState2 - Join Oppost	15%	37%	GrievanceState2 - Join Oppost
GrievanceState3 - Nonviolent	0%	39%	GrievanceState3 - Nonviolent
GrievanceState4 - Fight-Rebel	0%	17%	GrievanceState4 - Fight-Rebel
TOTAL	100%	100%	TOTAL

4.1) Correspondence Test

The correspondence test is whether the overall parameterization for the GSP tree-guided PMFserv agents in the bridge with the Civil Violence population will faithfully mimic the test data set. That is, by tuning the GSP trees of 1 leader and 160 villagers, and by connecting all that to the Civil Violence mode of spreading news and grievances, do we wind up with a simulation that seems to correspond to what happened in the real world test dataset? Specifically, we are interested in testing the null hypothesis that there is no statistically significant

correlation between real decisions and the simulated decisions. That is to say that real incidents and simulated base case are mutually independent.

When the simulation is run, one observes Leader of X trying some assistance measures initially (usually offering to set up Buddhist school and institutions) but maintaining a high police presence, and turning increasingly suppressive as the run proceeds -- Suppressing by Increasing Militarization and by Increasing Violence Unleashed. By the end of the run, the right side of Table 4 shows the emergence of a majority of the population resisting and fighting (non-violent as well as violent) against X. Specifically, it shows what percent of the population has been shifted from Neutral Grievance to higher states (recall the scale of earlier Section 3): GS0 (neutral) through GS4 (fight back). From the first graph, it can be seen that at the start, most villagers are near neutral and occupy GS0 and GS1, while a small percent start in GS2. The occupancy in lower grievance states fall with time, while that in higher grievance states climb. From about week 50 onwards, there is a fairly stable, though regularly punctuated equilibrium in which the highest occupied states are GS3 and GS4. This is an indication of progressive escalation of violence in the society since these two states represent a shift to fighting.

In order to compare this simulated grievance to that of the real world, we need some reliable measures of the population's grievance during actual events. Unfortunately, there are no survey or attitude results available. In the real world (test) dataset, the incident data was available, however, with a record of fatalities and injuries. There are a number of schemes for weighting those (e.g., depression and morale loss, lost income, utility metrics, others), however, here we take the simple approach of using weighted average of fatalities and injuries, where injuries are simply counted ($w=1$), but the weight on fatalities is 100. $IncidentSeverity = w_f \times fatalities + w_i \times injuries$. The result is a widely used proxy of how severe these incidents were: e.g., see Collier & Hoeffler (2001).. While severity is only an indirect measure of how the population might have felt, it is a measure that can be tested for correlation to the rise and fall of grievance expression due to leader actions in our simulated world.

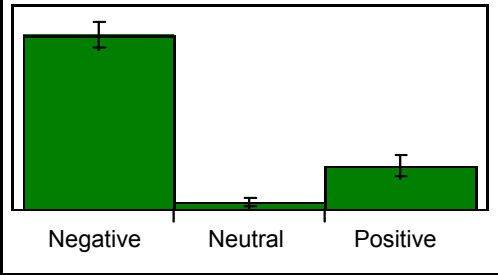
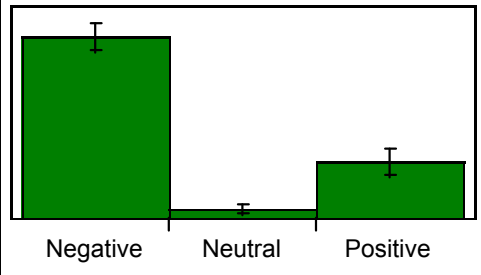
To conduct the comparison, we apply the non-parametric Kendall's Tau measure of correlation. This statistic estimates the excess of concordant over discordant pairs of data, adjusted for tied pairs. With a two sided test, considering the possibility of concordance or discordance (akin to positive or negative correlation), we can conclude that there is a statistically significant lack of dependence between base case simulation and observed grievances rankings at a confidence interval of 88%. Since there is a probabilistic outcome determining if a simulated leader's action choice will result in injury and fatality incidents (and how the news of these events are propagated through the cellular automata is probabilistic as well), we repeated the simulation runs thirty times and the confidence interval mentioned above is the mean across those 30 correlations. In sum, the null hypothesis is rejected and real (test interval) incident data and simulation results are related.

As to the leadership, we have detailed data and can conduct a correspondence test. Specifically, in the test dataset, the real world leader made 52 decisions affecting the population and that we sorted into positive, neutral, and negative actions. In the simulated world, LX made 56 action decisions in this same interval. The list of available actions was presented in Part I (Figure 1) and repeated in Section 2.1 above. At this level of classification (positive, neutral, negative), we were able to calculate a mutual information or mutual entropy (M) statistic between the real and simulated base cases (see Figure 3). M ranges from 0 to 1.0, with the latter indicating no correlation between

two event sets X and Y. Applying this metric, the mutual entropy values were found to be less than 0.05, indicating correlation between real and simulated data. With an M metric, one cannot make statements about the confidence interval of the correlation, however, the Leader in the current scenario seems equally faithful to his real world counterpart.

Figure 3 - Correlation of Simulated Leader vs. Real Action Decisions

Comparison of distributions to see Mutual Entropy (M). Reject H0 & Accept H1 if $M < 0.1$

	PMFserv-Simulated Leader's Actions		Real Leader's Chosen Actions	
<u>Distributions</u>				
<u>Mutual Entropy Calculations</u>	Joint Entropy of Sim & Real	1.396	$H(\text{SIM}, \text{REAL}) = - \sum p(\text{sim}_i, \text{real}_i) \log p(\text{sim}_i, \text{real}_i)$	
	Entropy of Sim	0.681	$H(\text{SIM}) = - \sum p(\text{sim})_i \log p(\text{sim})$	
	Entropy of Real	0.760	$H(\text{REAL}) = - \sum p(\text{real})_i \log p(\text{real})$	
	Mutual Entropy of Sim & Real	0.045	$M(\text{SIM: REAL}) = H(\text{SIM}) - H(\text{SIM} \text{REAL})$	

4.2) Turing Test and Experiment #2: Elasticity of Demand for Civil Rights

The Turing test looks beyond population statistics, and examines what transpires inside the heads of the various types of agents in the simulated world. In terms of the followers, the prior section shows them shifting from cooperate to fight. We are curious about the migration decision and how it unfolds. This was revealed earlier in Part I, Figure 3 which depicts a Radical Muslim Follower at the precise moment of his shifting from reluctant cooperation to resistance. As a fringe member of Y, he had a strong potential to radicalize and shift factions. Prior to that shift, his negative emotion bars were equally activated, but his positive ones were suppressed. This was due to his depression and anger over his perceived lowered VID of his group Y and how Leader X was mistreating it. Once he shifted the strength of his membership to resistant Group Z (separatists), his positive emotions were activated as shown in Figure 3 (Part I). In the Hirshman (1970) model of 'Loyalty, Voice, and Exit', this is also the moment he shifts from voice to exit.

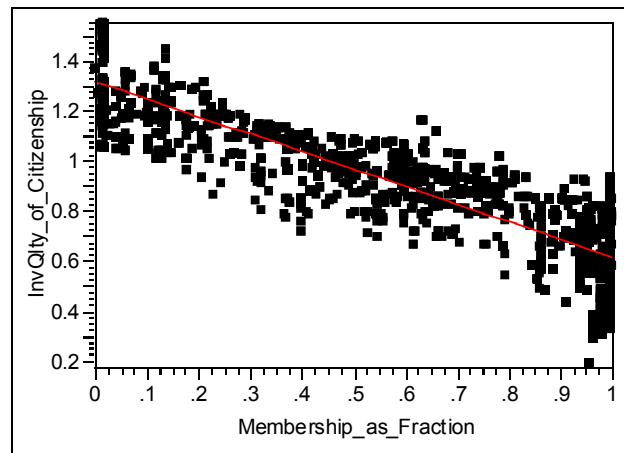
Likewise, one can inspect the GSP trees inside the head of the Leader (LX) and observe his emotional utility calculations for alternative action choices available to him on each tick. Structurally, these trees are identical to the GSP trees shown for the two Iraqi leaders in Figure 2 (Part I). His Bayesian weights were calibrated from frequency (prior odds) of action choices and tendencies in the training set. These comprise the 'base case'

personality of this leader. We then altered several of his key standards for treating others (outgroup are targets, sensitivity to life, scope of doing good, etc.) by 15% in either direction. Reducing these is equivalent to what the SMEs in the Iraqi case study attempted when they had another leader call and try to convince him to be more reasonable and tolerant. Raising these up by 15% is what might happen if LX grew more autocratic or if our prior odds are off a bit.

Since the Leader X's attributes lead directly to shifts in his course of action selections, these three versions of the leader were run to set up a range of potential futures for the followers. Figure 4 plots these along the Y-axis, or at least the Follower's reactions to these action sets. Thus, the y-axis represents increasing losses of civil rights or the Inverse Quality of Citizenship (InvQtyCitizenship) as measured by the Follower Group Y's calculated grievances (or VID from Part I of this article). The x-axis shows the decision of these agents to retain membership in Faction Y (these are the members of GS0, GS1, and GS2). Agents who leave and join Faction Z (GS3 and GS4) are not appearing in this plot. The plot thus shows that as long as conditions are not too intolerable, the entire population cooperates and remains in Faction Y. As conditions worsen, more and more agents exit and Faction Y's membership shrinks. This is what Hirshman refers to as the demand curve for civil rights. In FactionSim, we are able to fit the following linear regression to this demand curve with an Rsquare of 0.79

$$\text{InvQtyCitizenship} = 1.35364 - 0.8269131 (\text{Membership_as_Fraction}) \quad [2]$$

Figure 4 – Derived Demand Curve for Civil Rights by Faction Y's Followers



This curve in Figure 4 is derived from synthetic agents, however, it seems to describe a reality that the Leader of X fails to comprehend at his own political peril. In decades past, the rural Muslim villagers were well-behaved citizens of X. However, there is a new generation of young males who are willing to stand up for the civil rights and who are highly influential across the populace – hence the demand curve has a constant slope for most of its length. The Buddhist leader's ingroup bias, financial wealth, narrow scope of helping only his own faction to the north, and willingness to use violent repression seem to combine in the real world (and in our model of LX) and make him unable to comprehend this new reality. In the summer of 2005, LX had to impose martial law on these

1
2
3
4 provinces to try and quell the separatist movement. In the summer of 2006, with the approval of the monarch, a
5 military junta removed LX from power due to his mismanagement of this situation and economic issues.
6

7 Unlike the agents that are designed to try and win Prisoner's Dilemma type games, real world counterparts
8 are often less than rational, may be biased, and may have moral or other agendas. Cases like this one allow us to
9 compare the true behavior to what the theoretically optimal one might be. Some leaders demonstrate keen
10 understanding of the game, and purposely impose repressive measures to stifle voice, while opening borders to
11 promote exit. In this case, there were no borders to open, and LX simply tried to win by fighting. The rival leaders in
12 his own faction saw his tumbling political support before he did and removed him. His myopia cost him victory in
13 two games -- one against an allied faction's followers (Y) and the other against rivals in his own faction (X). We did
14 not set up this second game in the current case study, but we cautioned at the outset that a problem with game theory
15 is that agents often are involved in multiple games at once, games it is hard to even know about. We believe,
16 however, that the framework presented here allows one to set up and play out the larger scenario surrounding Leader
17 X. Our results to date on specific games we did set up give us confidence in the validity of this approach as we move
18 in the direction of having the agents try to manage multiple games at once – not as optimal ideals, but as realistic
19 counterparts of the true individuals.
20
21
22
23
24
25
26
27
28
29
30
31

32 **5) Lessons Learned and Next Steps**

33 The primary argument against rational game theory is its poor track record of prediction in matters of real
34 world conflict primarily because it often simplifies the game and agents to the point that they bear little resemblance
35 to the real world. For example, the normative prediction at the end of Section 2.2 was that rational agents (with
36 memory) in IPD games will find equilibrium in mutual cooperation. In contrast our Iraqi agents were far from
37 normative – at times attacking already defeated opponents who no longer had any resources to loot (perhaps to gain
38 political favor), at other times driven by hatred to suicidal attacks against overwhelmingly larger forces..
39 Nevertheless, game theory can be of help in structuring analysis. In Section 2 we invested effort in setting up dyadic
40 and triadic versions of the Iterated Prisoners' Dilemma game both to highlight where it can help as well as where it
41 falls short and needs descriptive agents. This lead us to 6 scenarios of the triadic game, each with predicted
42 equilibrium conditions for the one shot game. The scenario that appeared in Section 3 was S3.4 (a coalition fighting
43 against an aggressor) while Section 4 explored the pace at which agents shift from cooperating to fighting against an
44 aggressor (S2.2—> S2.1).
45
46
47
48
49
50

51 The primary argument against Behavioral Game Theory (BGT) in the social sciences is that there are few
52 first principles that all social scientists agree upon, the field is not mature. Still, that is no excuse for modelers to
53 “make up” their own rules and algorithm for how groups behave, nor is it justification to just create rational actors.
54 The alternative we explored here is the systems approach where we take game elements and agent relations and
55 cognition and break these into a system of components (sub-systems). Each component has encapsulated
56 functionality, preserves inter-relationships between components, and applies domain theories/knowledge to keep
57
58
59
60
61
62
63
64
65

1
2
3
4 realism in the components to the extent possible. This is the systems approach and it can be applied recursively to
5 any component. An advantage of this approach is to encapsulate behavior so components can be modeled at varying
6 resolution without affecting how the collection of components interact. An example of this was the (economy,
7 security, or political support) resource tanks that we currently model as stacks of poker chips that grow or fall. One
8 can plug in finer resolution models for any given tank without affecting overall system performance.
9

10
11 Another example of this concerns the parameters internal to a given agent where we try and synthesize
12 best-of-breed and well-respected social science models for leadership, group dynamics, and the hearts and minds of
13 the populace. Once again, the systems approach prevails and these are in 6 cognitive components with inter-relations
14 between them nicely externalized as explained in Part I of this paper. If an analyst dislikes some of these, they can
15 readily be over-ridden. Where they are apropos, such models reduce the dimensionality to the traits and factors they
16 require, and where these are applied, we can use training datasets, fill in the traits and factors of archetypical as well
17 as real characters, conduct validation tests, and treat these parameters as no longer existent. That is they are no
18 longer independent variables clouding the larger DIME-PMESII analyses, but are swept out of the way by first
19 principles, training data, and validity tests before DIME-PMESII studies even begin.
20
21

22 An argument against the realism and richness of this approach is the 'curse of dimensionality' -- the
23 explosion of parameters that demand unattainable amounts of training datasets. This happens in social science
24 problems when one tries to drag ever more parameters in to try and explain variances and fluctuations in the world
25 being modeled. In this case the crime of over-fitting often occurs, where the model has so many variables tuned to
26 one data set, but it can't then successfully explain a different dataset. The systems approach, however, provides the
27 dual benefit of synthesizing the social science models into a wholism at the same time that it uses them as domain
28 knowledge to remove these parameters from the frame for DIME-PMESII analysis. Even if we hypothesize two
29 versions of a given leader (e.g., one more benevolent the other more autocratic), the hundreds of parameters inside
30 them are reduced down to just those two. Unlike the evolutionary tradition where personas must be mutated, this
31 approach of profiling real personalities allows one to watch what they do and learn – what behavior emerges from
32 current actors. Normative or rational actors can of course be scripted to have diverse payoff functions and action
33 preferences (e.g., normative altruist, autocrat, grim, etc.). However, since descriptive agents are not scripted, but are
34 personality profiled, one is freed of the need to guess which normative script to have them follow. Their action
35 choices emerge dynamically as the game unfolds.
36
37

38 It is worth dwelling a bit on the benefits that were observed and the lessons learned from the case studies of
39 this paper. For one thing, the descriptive agents passed validity assessment tests in both conflict scenarios—the Iraqi
40 leader agents were passed after extensive SME evaluation and the SE Asia leader and followers passed separate
41 correspondence tests (correlations of over 79%). Validity is a difficult thing to claim, and one can always devise
42 new tests. A strong test, however, is the out-of-sample tests that these agents also passed. Thus the SE Asian leader
43 and followers were trained on different data than they were tested against (see Sect.4). Further, the complete
44 structure of the model of the leaders was originally derived in earlier studies of the ancient Crusades (Silverman et
45 al. 2005) and this was transferred to the SE Asian and Iraqi domains. The only thing updated was the values of the
46 weights for GSP trees and various other group relations and membership parameters – derived from open sources.
47
48
49
50
51
52
53
54
55
56
57
58
59
60
61
62
63
64
65

1
2
3
4 So the structure of the leader model also survived and passed two out-of-sample tests relative to the Crusades
5 dataset. While these may not be the ultimate tests, they are sufficient for our purposes and in order to consider the
6 descriptive agents to be components that reduce the dimensionality issues.
7

8
9 As mentioned in Sect.1, a major objective of FactionSim is to support experiments on synthetic agents to
10 identify those policy instruments that will most influence the real-world agents they represent – ie, EBO studies. In
11 terms of the experiments attempted and presented, we showed 3 Iraqi runs (mean of 100 trials each) as well as 3
12 Country-T runs. The principal independent variable evaluated in the illustrative Iraqi experiment concerns how
13 much outside support is reaching the two protagonists – CentralGov and Insurgents. When CentralGov is heavily
14 supported and the Insurgents less so, the fighting continues throughout the 2 year run. When only the Insurgents are
15 supported, the CentralGov fails, and when the borders are fully closed and no group receives outside support, the
16 insurgency ultimately fails. CentralGov is fairly benevolent to the Shia's in all runs but in this closed-borders run
17 they begin to reach out to the Sunnis as well. These runs suggest the elasticity of conflict with respect to outside
18 support is positive, and with no interference, the country seems able to right itself, although we in the West might
19 not like the outcome. Of course these runs only include 4 of the many factions one could set up and run, plus due to
20 page limits, we only displayed the effects of actions upon the Security Tank, and not other resources of the factions.
21 Also, while the leaders' internal parameters are not independent variables in this set of runs, one can 'open a
22 window that reveals what is driving their emotions, grievances, relationships, and decision choices, some of which
23 was shown in Part I.
24

25
26 The other experiment presented (Sect.4) concerned the internal parameters of the leader of a SE Asian
27 nation. Specifically, we were interested in the elasticity of follower cooperation as the leader's behavior shifted – a
28 demand curve for civil rights. Or, put another way, to what extent would followers (Muslim moderates and radicals)
29 exit and join an insurgency as the Buddhist leader's policies became more draconian. Our population model
30 involved a cellular automata with 1,360 initially neutral agents influenced in their neighborhoods and schools by 160
31 PMFserv agents, half of which were moderate, half radical. Our independent variable in the experiment was the loss
32 of civil rights of the populace, and we found it regressed directly with membership loss (and insurgency growth)
33 with an R^2 of .79. Once again, Part I shows one of the radical followers's emotions at the moment of exiting. This
34 experiment is interesting since it predicts loss of control of the populace, a reality that occurred a year after our runs
35 when the leader declared martial law. He subsequently was removed from office over this affair.
36

37
38 In summing up, the two experiments illustrate the value of descriptive agents for extending game theory.
39 The entire point of insisting on well-respected models inside and on validation efforts for the descriptive agents is so
40 one can have trust that BGT experiments on these agents will yield insights about the alternative policies that
41 influence them. Both experiments serve to illustrate how analysts might use BGT and FactionSim to support their
42 tradecraft, to explore policy alternatives and robustness, and to identify parameter elasticities or sensitivities. Hence
43 an important line of investigation in our future work will be to develop a range of more systematic and less effort-
44 intensive statistical techniques that can be used by practitioners as preliminary steps in the construction and
45 evaluation of policy alternatives. By linking such parameters to specific policy instruments, practitioners may then
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60
61
62
63
64
65

1
2
3
4 be able to identify those policy alternatives that are potentially most effective in guiding agent behavior toward
5 desired outcomes.
6

7
8
9 **ACKNOWLEDGEMENT:** This research was partially supported by the US Government (game engine, leader
10 models), DARPA (follower models, case study frameworks), AFOSR (correspondence tests, all experiments,
11 FactionSim), and by the Beck Scholarship Fund. Lockheed Martin/ATL prepared an early version of the cellular
12 automata framework we updated and used here. Also, we thank the US Government for their guidance and
13 encouragement, and many students for their help, though no one except the authors is responsible for any statements
14 or errors in this manuscript.
15
16
17

18 19 **REFERENCES** 20

21
22 Armstrong, JS (2002). Assessing game theory, role playing and unaided judgment. *International Journal of*
23 *Forecasting*, 18, 345-352.
24

25
26
27 Axelrod, R., and Bennett, S. (1993) A Landscape Theory of Aggregation, *British Journal of Political Science* 1993
28

29
30 Camerer, C (2003), *Behavioral Game Theory*, Princeton: Princeton Univ. Press.
31

32
33 Collier, P, Hoeffler, A, (2001), "Greed and Grievance in Civil War", Washington DC: World Bank, avail. at
34 www.worldbank.org/research/conflict/papers/greedandgrievance.htm
35

36
37
38 De Marchi, S (2005). *Computational and Mathematical Modeling in the Social Sciences*, Cambridge
39

40
41 Dutta, P. K. (2000) *Strategies and Games: Theory and Practice*. Cambridge, MA: MIT Press.
42

43
44 Epstein, J., Steinbruner, JD, Parker, MT, (2001), "Modeling Civil Violence: An Agent-Based Computational
45 Approach," *Proceedings of the National Academy of Sciences*. Washington DC: Brookings (CSED WP#20).
46

47
48 Field, A. J. (2001). Altruistically inclined? The behavioral sciences, evolutionary theory, and the origin of
49 reciprocity. In T. Kuran (Ed.), *Economics, Cognition, and Society*. Ann Arbor: University of Michigan Press.
50

51
52
53 Finus, M (2002), "New Developments in Coalition Theory," in Rauscher, M (ed) *The International Dimension of*
54 *Environmental Policy*, Dordrecht: Kluwer
55

56
57
58 Friedman, J. (1971). A non-cooperative equilibrium for supergames, *Review of Economic Studies* 38, 1-12.
59
60
61
62
63
64
65

Green, K. C. (2002). Forecasting decisions in conflict situations: a comparison of game theory, role playing and unaided judgment. *International Journal of Forecasting*, 18, 321–344

Heuer, R. J., Jr. (1999). *Psychology of Intelligence Analysis*. Washington, DC: Center for the Study of Intelligence, Central Intelligence Agency.

Hirschman, A.O. (1970). *Exit, voice, and loyalty*. Cambridge, MA: Harvard University Press.

Johns, M., *Deception and Trust in Complex Semi-Competitive Environments* (Doctoral dissertation, University of Pennsylvania, 2006)

Kaneko, M. (1982). Some remarks on the folk theorem in game theory, *Mathematical Social Sciences*, Elsevier, vol. 3(3), pages 281-290, October.

McCrabb, MJ, Caroli, JA (2002), “Behavioral Modeling and Wargaming for Effects-Based Operations,” *Proc. Military Operations Research Society Annual Meeting*, Washington DC: MORS.

Macy, MW, Flache, A, (2002) “Learning Dynamics in Social Dilemmas,” *PNAS*, May 14, v. 99, Sup 3, 7229-7236

Sageman, M. (2005), *Understanding Terror Networks*, Philadelphia: U. of Pennsylvania Press

Silverman, BG, Bharathy, G. (2005), “Modeling the Personality & Cognition of Leaders,” in 14th Conf on Behavioral Representations In Modeling and Simulation, SISO (www.sisostds.org), May.

Silverman, B. G., Johns, M., Cornwell, J., & O’Brien, K. (2006a). Human Behavior Models for Agents in Simulators and Games: Part I – Enabling Science with PMFserv. *Presence*, v. 15: 2, April.

Silverman, B. G., O’Brien, K., Cornwell, J. (2006b). Human Behavior Models for Agents in Simulators and Games: Part II – Gamebot Engineering with PMFserv. *Presence*, v. 15: 2, April.

Silverman, BG, Bharathy, G., Nye, B (2007a), “Gaming and Simulating EthnoPolitical Conflicts” in Proc. Descartes Conf on Mathematical Modeling for Counter-Terrorism (DCMMC), NYC: Springer,

Silverman, BG, Bharathy, GK, Johns, et al. (2007b), “Socio-Cultural Games for Training and Analysis”, (submitted for publication) avail at: www.seas.upenn.edu/~barryg/CultureGames.pdf

1
2
3
4 Simari, G. I., & Parsons, S. (2004). On approximating the best decision for an autonomous agent. *3rd AAMAS Conf.*,
5 W16: Workshop on Game Theory and Decision Theory. NYC
6
7

8
9 Wood, E. J (2003), "Distributional Settlements and Civil War Resolution: Stakes, Expectations, and Optimal
10 Agreements," under revision for re-submission to *The Journal of Conflict Resolution*.
11
12
13
14
15
16
17
18
19
20
21
22
23
24
25
26
27
28
29
30
31
32
33
34
35
36
37
38
39
40
41
42
43
44
45
46
47
48
49
50
51
52
53
54
55
56
57
58
59
60
61
62
63
64
65

Glossary

Terms	Description
S2	Pertains to dyadic scenarios, can be considered a simplified subgame in a triadic interaction. Dyadic scenarios are described without S2 prefix.
S3	Pertains to triadic scenarios.
S3.1, S3.2...S3.6	Each one is a triadic scenario.
S2x[FxFy]	Payoff to x in a dyadic scenario, when Both x and y are fighting. Mutual conflict
S2x[FxCy]	Payoff to x in a dyadic scenario, when x is fighting while y has compromised.
S2x[CxFy]	Payoff to x in a dyadic scenario, when y is fighting while x has compromised.
S2x[CxCy]	Payoff to x in a dyadic scenario, when both x and y have compromised. Mutual compromise.
S3x[FxFy, FxFz, FyFz]	Payoff to x in a triadic scenario, when x, y and z are fighting with each other. Mutual conflict
S3x[CxFy, CxFz, CyCz]	Payoff to x in a triadic scenario, when the aggressors y and z independently attack a passive x
S3x[CxCy, CxFz, CyFz]	Payoff to x in a triadic scenario, when z attacks coalition of x and y, who do not fight back
S3x[CxCy, FxFz, FyFz]	Payoff to x in a triadic scenario, when z is fighting with coalition of x and y
S3x[CxCy, CxCz, CyCz]	Payoff to x in a triadic scenario, when there is mutual cooperation/ compromise
i	Discount rate discounting future payoffs to account for time value of payoffs
X, Y, Z..	Leaders in the world. Also used as x,y,z when subscripted
Qj	Level of attack j
Qjzx	Level of attack that denotes the attack is by leader Z on leader X
Qjz_xy	Level of attack where the attack is by leader Z on the coalition of leader X and Y
QjZY_X	Level of attack which denotes that the attack is by the coalition of leaders Z and Y on leader X
Rx, Ry, Rz	Total resources of X, Y, Z
R2	The total resources in a dyadic interaction $R_x + R_y = R_2$
R3	The total resources in triadic interaction be $R_x + R_y + R_z = R_3$
Rdy	Disputed or contested Resource share that belongs to Leader y when both x and y are compromising
Rdx	Disputed or contested Resource share that belongs to Leader x when both x and y are compromising
Rd	Total pool Disputed or contested Resource that will be shared by the Leaders, when when both x and y are compromising
ΔK_{xy}	Changed in dyadic relationships between x and y. This is a function of relationships between

	the leaders as well as the actions taken.
CstB	The cost of staging a battle in a dyadic interaction
P _x	Probability of winning in a battle, and is proportional to level (effort) of attack (Q _{jyx}) and relative strength $R_y/(R_x+R_y)$ of the attacker => $P_x = (Q_{jyx} \cdot R_y)/(R_x+R_y)$.
Q _{jyx} ($R_y/(R_x+R_y)$) R _{dx}	The expected loss in a given battle for a target is proportional to the level of attack, likelihood of success and the level of resource contested. This is const.(relative strength of attacker)(contested resource of attacked)
Q _{jzx} (R_z/R_3)R _{dx}	Expected losses to x due to being attacked by z using relative resources available (R_z/R_3). The attack takes place on the contested resource R _{dx} , which belongs to x.
emV(F _x , C _y)	Emotional payoff (non-material utility) for X from X fighting while Y compromising
S _{._.x} (t)	Refers to the payoff for x in scenario S _{._.} occurring in time step t.

Barry Silverman is professor of Electrical and Systems Engineering, Medicine, and Wharton/OPIM at the University of Pennsylvania as well as a fellow of IEEE, AAAS, and Washington Academy of Science, and honoree of several professional societies for best paper, research, and teaching awards. He has authored or edited 12 books, 125 articles, 1 board game, and 7 copyrighted software systems on intelligent agents, knowledge management and virtual worlds for enhancing human performance and adaptivity.

Gnana K. Bharathy recently completed his doctoral work in Systems Engineering. During the course of his dissertation work, Gnana has developed a systems methodology for integrating social system frameworks and modeling human behavior through knowledge engineering based process, and has employed the same to create several models of leaders and followers in situations involving conflict-cooperation. His dissertation has recently been awarded the INCOSE-Stevens award for promising research in systems engineering and integration

Ben Nye is currently pursuing his doctorate in Electrical and Systems Engineering, researching alongside Barry Silverman. Current research work he is involved in focuses on simulating decision making in sets of hierarchal groups, working towards automated exploration of the state space. In addition to a B.S. in Computer Engineering, he also has a significant background in psychology.

Tony Smith is a Professor of Systems Engineering and Regional Science at the University of Pennsylvania. His primary area of research is in the theory and application of probabilistic models to spatial interaction behavior. A secondary area of research is in transportation and land use modeling. Dr. Smith holds his PhD from the University of Pennsylvania. He has written two textbooks and numerous journal articles. Dr. Smith was designated the 1999 recipient of the Walter Isard Award for Distinguished Scholarship by the Regional Science Association International.