



American Finance Association

Corporate Board Composition, Protocols, and Voting Behavior: Experimental Evidence

Author(s): Ann B. Gillette, Thomas H. Noe, Michael J. Rebello

Source: *The Journal of Finance*, Vol. 58, No. 5 (Oct., 2003), pp. 1997-2031

Published by: Blackwell Publishing for the American Finance Association

Stable URL: <http://www.jstor.org/stable/3648181>

Accessed: 24/11/2009 09:24

Your use of the JSTOR archive indicates your acceptance of JSTOR's Terms and Conditions of Use, available at <http://www.jstor.org/page/info/about/policies/terms.jsp>. JSTOR's Terms and Conditions of Use provides, in part, that unless you have obtained prior permission, you may not download an entire issue of a journal or multiple copies of articles, and you may use content in the JSTOR archive only for your personal, non-commercial use.

Please contact the publisher regarding any further use of this work. Publisher contact information may be obtained at <http://www.jstor.org/action/showPublisher?publisherCode=black>.

Each copy of any part of a JSTOR transmission must contain the same copyright notice that appears on the screen or printed page of such transmission.

JSTOR is a not-for-profit service that helps scholars, researchers, and students discover, use, and build upon a wide range of content in a trusted digital archive. We use information technology and tools to increase productivity and facilitate new forms of scholarship. For more information about JSTOR, please contact support@jstor.org.



Blackwell Publishing and American Finance Association are collaborating with JSTOR to digitize, preserve and extend access to *The Journal of Finance*.

<http://www.jstor.org>

Corporate Board Composition, Protocols, and Voting Behavior: Experimental Evidence

ANN B. GILLETTE, THOMAS H. NOE, and MICHAEL J. REBELLO*

ABSTRACT

We examine voting by a board designed to mitigate conflicts of interest between privately informed insiders and owners. Our model demonstrates that, as argued by researchers and the business press, boards with a majority of trustworthy but uninformed “watchdogs” can implement institutionally preferred policies. Our laboratory experiments strongly support this conclusion. Our model also highlights the necessity of penalties on insiders when there is dissension among board members. However, penalties for dissent appeared to have little impact on the experimental outcomes.

CONSIDER THE SITUATION OF OWNERS of a corporation when they entrust the fate of their institution to groups of insiders. Because the owner-preferred allocation may be contingent on private information possessed by the insiders, owners need a mechanism to mitigate conflicts of interest with insiders. In practice, owners entrust the governance of corporations and other institutions to boards and committees consisting of a mix of insiders and outsiders. Corporate boards are increasingly being dominated by independent outside directors (see, e.g., Spencer Stuart (1997)), a trend mirroring the recommendations of the National Association of Corporate Directors and The Business Roundtable.¹

The advocacy of independent outsider-dominated boards is surprising as research has produced weak or mixed results on the effectiveness of outsiders on boards. For example, Weisbach (1988) finds that, when there is a majority of outside directors, CEO turnover is greater. However, the likelihood of CEO replace-

* Gillette and Rebello are on the faculty of Georgia State University, and Noe is on the faculty of Tulane University. We would like to thank Rick Green (the editor), an anonymous referee, participants at the Atlanta Finance Workshop, the Caltech Bray Seminar, College of William and Mary, the Economic Science Association Meetings, the 2000 Financial Management Association Meetings, Kennesaw State University, the SIRIF Behavioral Finance Conference, U. Cattolica del Sacro Cuore, and the 2001 Western Finance Association Meetings. We also wish to thank Lucy Ackert, Christopher Anderson, and Chuck Schnitzlein for useful comments. Gillette acknowledges research support from the Robinson College of Business, Georgia State University and The Federal Reserve Bank of Atlanta. We are also grateful to Ping Hu, Bing-Xuan Lin, and Gwendolyn Pennywell for research help. Any errors are our own.

¹See, for example, the Report of the National Association of Corporate Directors Blue Ribbon Commission on Director Professionalism (1996) and The Business Roundtable Statement on Corporate Governance (1997).

ment is only slightly higher for these firms. Mikkelsen and Partch (1997) also find little evidence of a relationship between CEO tenure and board composition. In a corporate control context, Byrd and Hickman (1992) document a more favorable response to acquirers with majority-outsider boards, while Subrahmanyam, Rangan, and Rosenstein (1997) find the opposite tendency for bank acquisitions. Studies of the direct relationship between board composition and firm performance have also produced mixed results. For example, while Baysinger and Butler (1985) document a positive relation between outsider membership on boards and return on equity relative to industry, Yermack (1996) reports a negative relation between the proportion of outside directors and Tobin's q . Further, a number of other studies find no significant relation between the mix of directors and same-year firm performance (see, e.g., Hermalin and Weisbach (1991) and Mehran (1995)).

A lack of clear evidence for the effectiveness of outside directors has fueled the continuing debate on the effectiveness of outsider-majority boards. This lack of support for the effectiveness of outsider-majority boards could be the result of impediments to empirical research pointed out by a number of authors, including Hermalin and Weisbach (1991), Bhagat and Black (1999), and MacAvoy and Millstein (1999). These impediments include difficulties in measuring the day-to-day effect of board composition on corporate performance, poor disclosure regarding consulting opportunities and grants to directors' employing institutions, the omission of outside directors' compensation-related incentives from studies, and flawed proxies for the level of board independence.

We examine board effectiveness using a technique that enables us to address many of the confounding problems faced by the empirical research on this subject—laboratory experiments with human subjects. Our subjects play a game that embeds several features of real life boards. A board is constituted of two groups of subjects—insiders and independent outsiders (watchdogs), with watchdogs commanding the majority on the board. Together, they must decide whether to accept a project. The project's fate is decided by majority vote. The institutionally preferred policy calls for the project to be accepted if it is value increasing, and rejected otherwise.

Watchdogs are unable to discriminate between value-increasing and value-decreasing projects whereas insiders have private information that enables them to distinguish between the two types of projects. To focus the analysis on the central task facing boards, mitigating conflicts of interest between owners and insiders, the incentive structure for each group of subjects is chosen to ensure that watchdog incentives are aligned with those of the firm's owners while insiders' incentives are misaligned. This misalignment arises because insiders obtain private benefits from investment even if it destroys firm value.

Insiders are subject to penalties resulting from board dissent. Such penalties are common in real-world institutions (Warther (1998)), a prime example being the dismissal of the former CFO and board member of Apple Computer Inc., Joseph Graziano. He was "shot-down" by the board after placing strong objections to the CEO's yearly business plan, and resigned the next day. In the context of our game, this aspect of insider incentives is captured by a positive probability that insiders are penalized following a split vote (some yes, some no).

Subject behavior in the experiments is compared to the predictions of a model. The model has both “efficient” and “inefficient” equilibrium outcomes. Under the efficient outcome, which is supported by many equilibria, the institutionally preferred policy is adopted and insider votes are not split. Because all insiders receive the same information regarding project quality, if a single insider cannot sway the outcome of the decision with his recommendation, the penalty for dissent ensures that it is in his interest to go along with the other insiders. Thus, if all other insiders are acting in the firm’s interest, he will also do so. The same logic can lead insiders to choose self-serving policies. Watchdogs, though uninformed, have rational expectations; thus, they correctly conjecture the probability distribution for project choice that results if insiders were allowed to determine project choice. As the watchdogs command the majority of votes on the board and have an incentive to veto egregious policy choices, in equilibrium, insiders support the institutionally preferred policy (Palfrey (1992) and Palfrey and Srivastava (1993)).

However, a watchdog-dominated board may not be able to ensure that the institutionally preferred policy is undertaken. First, coordination between agents may fail and convergence to Nash equilibrium may not occur. Second, coordination between insiders may be “too successful,” leaving equilibria supporting the institutionally preferred outcome susceptible to the resulting coalition of insiders. In other words, equilibria may not be “coalition-proof” (Bernheim, Peleg, and Whinston (1987)): There may exist “self-enforcing” strategy vectors involving simultaneous deviations by a subset of agents that produce a higher payoff to that subset.²

Only equilibria supporting the inefficient outcome, where the project is always rejected regardless of its quality, are coalition proof. In these equilibria, watchdogs always block acceptance of the project, preventing a coalition of insiders from destroying firm value via investments in poor projects. Although the inefficient outcome is strictly dominated by the efficient outcome in the sense that all agents’ payoffs are strictly higher under the efficient outcome, equilibria supporting the efficient outcome are not coalition proof. In fact, all coalition-proof equilibria produce the inefficient outcome.

Our experiments provide strong evidence that outsider-dominated boards implement institutionally preferred policies. The institutionally preferred outcome resulted the vast majority of times even in our central treatments that were designed to facilitate the formation of insider coalitions. Surprisingly, subjects voted for the institutionally preferred policy even when the experimental treatments did not include a penalty for the dissent. Consequently, the incidence of the institutionally preferred policy was not statistically different when the penalty for dissent was dropped from the experimental design.

Experiments designed to examine the robustness of these results provided further evidence that outsider dominated boards tend to produce institutionally

²In this context, self-enforcing strategy vectors are ones where, holding fixed the strategy set of the nondeviating agents, the deviating agents do not have an incentive to “double-cross” other deviating agents by defecting from the deviating coalition.

preferred outcomes. These robustness experiments allowed for a variety of communication structures between agents and were motivated by both recommendations for governance reform and research on communication between agents. For example, in 1994 the board of General Motors issued guidelines designed to increase the effectiveness of its board. These guidelines, which have since been championed by CalPERS, among others, included guidelines on communication between board members such as a requirement that independent directors meet without any insider present at least two or three times a year. Researchers too have argued that such changes in the structure of communication between agents could affect the effectiveness of mechanisms such as boards (Milgrom and Roberts (1996)) and the efficacy of cheap talk (Farrell and Rabin (1996)).

In addition to contributing to the literature on board effectiveness, our work is also related to a number of other streams of research. Within the experimental economics literature, our analysis lies at the interstices of research on communication, voting, and implementation. A number of researchers have shown that nonbinding preplay communication affects the outcomes of experiments (e.g., Cooper et al. (1989) and Forsythe, Lundholm, and Rietz (1999)). This communication literature led us to consider the effects of a wide variety of communication protocols. In contrast to these papers, however, we treat communication as a control variable rather than an object of study. Research on voting games includes Forsythe et al. (1996) and Rietz, Myerson, and Weber (1998). Unlike these papers, our voting game is characterized by information asymmetry regarding the relationship between votes and subject payoffs. The experimental literature on implementation focuses on experiments with two agents in games of complete information, such as Abreu-Matsushima mechanisms (Sefton and Yavas (1996)). In contrast, we consider implementation in settings characterized by many agents and incomplete information. Thus, unlike two-agent settings, our experimental design makes coalition formation, and the format in which communication takes place, central.

The remainder of the paper is organized as follows. Section I describes the game, delineates the equilibria, and analyzes their properties. Section II describes the experimental design. Section III presents our results. Section IV contains a discussion of the results from robustness experiments. The final section concludes the paper and suggests directions for future research. Proofs are confined to Appendix A.

I. Model

In this section, we present our model and highlight the basic tension inherent in the model—there is a unique coalition-proof outcome and a unique institutionally preferred outcome which can be supported by Nash equilibria. The institutionally preferred outcome is not coalition proof. However, all agents, even the insiders, are better off under the institutionally preferred outcome.

At an intuitive level, the model is straightforward. However, formalizing our argument is tedious for two reasons. First, representing communication over a

general message space requires copious notation. Second, the standard definition of coalition-proof Nash equilibria is inductive, and thus the machinery of mathematical induction must be utilized for all formal proofs relating to coalition proofness. In the interest of rigor, we provide a formal analysis of the equilibria below. However, to the extent possible, we have confined technical discussion and notation to Appendix A.

A. Agents and Information

The board consists of $[N] = \{1, 2, \dots, N\}$ agents. The first w are watchdogs who belong in the set $[w] = \{1, 2, \dots, w\}$. The remaining i agents are insiders from the set $[i] = \{w + 1, w + 2, \dots, N\}$. We assume that $w > i > 2$, ensuring that watchdogs have a voting majority. Before voting, insiders receive an information signal, s , revealing whether the project is good (G) or bad (B). Project acceptance is value increasing if the observed signal is G and destroys value when it is B . Watchdogs cannot observe project quality but believe that it is good (bad) with probability π ($1 - \pi$).

B. Actions and Strategies

The game consists of two stages: (1) a communication stage and (2) a decision stage. All communication is “cheap” in that it has no direct effect on agent welfare. It occurs after insiders receive their information signal.³ Let the i th watchdog’s communication strategy be represented by the message $\mathbf{m}_i^w \in M$, where M is a message space. Similarly let insider i ’s communication strategy following the observation of his private signal be given by $\mathbf{m}_i^I \in M^{(G,B)}$.

Conditioned on the messages exchanged in the first stage, the board votes on the project. Insiders either vote in favor of the project, $\mathcal{Y}$, or against it, $\mathcal{N}$. Watchdogs either vote against the project, $\mathcal{N}$, or abstain, $\mathcal{A}$. Let $v = (v_1, \dots, v_N)$ represent the vector of board votes, V represent the set of all possible vote vectors, v^W represent the subvector of watchdog votes, and let v^I represent the subvector of insider votes. For any vector (or subvector of) v , let $\#\mathcal{Y}(v)$ represent the number of yes votes and let $\#\mathcal{N}(v)$ represent the number of no votes.

The project is accepted if strictly more votes are cast for the project than against the project. We represent acceptance with an indicator function a , where $a(v) = 1$, if and only if $\#\mathcal{Y}(v) > \#\mathcal{N}(v)$. Voting exhibits consensus if no insider voted for an investment policy different from the policy adopted. Let $c: V \rightarrow \{0, 1\}$ be the indicator function for whether all insiders support the majority vote. That is, $c(v) = 1$, if and only if $\#\mathcal{Y}(v) > \#\mathcal{N}(v) \Rightarrow \#\mathcal{N}(v_I) = 0$ and $\#\mathcal{Y}(v) \leq \#\mathcal{N}(v) \Rightarrow \#\mathcal{Y}(v_I) = 0$. The imposition of the penalty following lack of consensus on the board is stochastic and is represented by a zero-one-valued ran-

³Theoretically, it does not matter whether the voting space is “yes,” “no,” and “abstain” for both types of traders, or the restricted voting space described above. Restricting the watchdogs’ space to “abstain” or “no” however, helps simplify the strategy space and thus reduces the complexity of the coordination problem.

dom variable $\tilde{z}$ that is independent of $\tilde{s}$. A penalty is imposed on all insiders when $z = 1$. The probability of this occurring is represented by ρ .

C. Payoffs

Agents' payoffs depend on four events: (1) the insiders' signal, $\tilde{s}$; (2) whether insider voting exhibited consensus; (3) whether the project was approved (A) or rejected (R); and (4) whether a penalty $P > 0$ is assessed on insiders following board dissent. Thus, we can write the ex post payoff of agent j which we represent by U_j as follows: If j is an insider then

$$U_j(v, \omega) \equiv U_I(v, \omega) \equiv a(v)x(I, A, s(\omega)) + (1 - a(v))x(I, R, s(\omega)) - P(1 - c(v))\tilde{z}(\omega) \quad (1)$$

If j is a watchdog then

$$U_j(v, \omega) \equiv U_W(v, \omega) \equiv a(v)x(W, A, s(\omega)) + (1 - a(v))x(W, R, s(\omega)) \quad (2)$$

The rankings of the payoffs, x , given in the above equations are as follows:

$$x(I, A, G) > x(I, R, G), x(I, A, B) > x(I, R, B); \\ x(W, A, G) > x(W, R, G), x(W, A, B) < x(W, R, B). \quad (3)$$

$$\pi x(W, R, G) + (1 - \pi)x(W, R, B) > \pi x(W, A, G) + (1 - \pi)x(W, A, B) \quad (4)$$

$$\rho x(I, A, s) - \rho P > x(I, R, s), s = G \text{ or } B \quad (5)$$

Assumption (3) expresses the fact that insiders prefer acceptance of the project regardless of its quality and watchdogs prefer to accept the project only when it is good, that is, $s = G$. Assumption (4) implies that, if watchdogs have to make their accept/reject decision based on their prior information, they prefer to reject the project. Assumption (5) implies that insiders are willing to pay the price for lack of consensus if, by paying this price, they will ensure acceptance of the project.

D. Results

First, we consider coalition-proof equilibria. As we show in Appendix A, coalition proofness ensures that insiders and watchdogs each act as if each group were a single agent. This restriction implies that insiders will force acceptance of the project whenever acceptance is not blocked by watchdogs. Knowing that insiders will not condition their support for the project on the information signal, watchdogs realize that they must choose between accepting the project for both information signals and rejecting it for both signals. By Assumption (4), watchdogs prefer blocking the project for both signals to accepting the project for both signals. Thus, watchdogs will always block the project. Realizing this fact,

insiders will also vote against the project, so as not to call down the penalty for lack of consensus. The next theorem summarizes these results.

THEOREM 1: *In all coalition-proof equilibria, (a) all insiders vote to reject the project under both information signals and (b) under both information signals, enough watchdogs vote against the project to ensure that, even if all insiders were to switch their votes to acceptance, the project would still be rejected. Moreover, coalition-proof equilibria exist in which all agents vote against the project under all information signals.*

Proof: See Appendix A.

In the coalition-proof equilibrium, watchdogs and insiders will not deviate to an equilibrium in which watchdogs abstain and the efficient project choice is implemented, because watchdogs realize that an agreement to switch to this equilibrium would be betrayed by insiders. Thus, watchdogs veto. Given the watchdog veto, insiders recognize that supporting the project is futile and thus, they vote against it.

Next we define efficient equilibria.

DEFINITION 1: *An outcome is efficient if (a) the project is accepted only when insiders observe signal G and rejected only when they observe signal B and (b) there is no chance that insiders will incur the penalty for a failure of consensus.*

We now show that Nash equilibria producing the efficient outcome exist. However, these equilibria are not coalition proof. A Nash equilibrium only requires that agents consider the effects of unilateral deviations from candidate equilibrium strategy vectors. Insiders know that if other insiders are providing “honest” recommendations by sending different messages depending on the information signal, unilateral efforts of a single insider to ensure project acceptance under both signals is futile and will simply call down the penalty for consensus failure. Thus, insiders will not deviate from the candidate equilibrium. Because the candidate equilibrium produces the highest possible payoff to watchdogs, they will not deviate from the equilibria implementing the efficient outcome.

THEOREM 2: *In all Nash equilibria that implement the efficient outcome, all insiders vote $\mathcal{Y}$ if they observe the signal G and vote $\mathcal{N}$ if they observe the signal B . Moreover a Nash equilibrium implementing the efficient outcome exists.*

The above result demonstrates that, as argued by commentators and researchers, outsider-dominated boards can implement institutionally preferred policies. However, Theorem 1 shows that the implementation of such policies is not guaranteed. Board failure is likely when insiders can coordinate their actions and act as a coalition. The experiments described in the following section are intended to provide insights into the functioning of boards by examining

subject behavior in settings characterized by the two equilibria described in this section.

II. Experimental Design

The set of experiments consisted of several treatments. Each treatment involved between one and three experimental sessions and lasted 10 rounds. The experimental subjects were undergraduate finance majors and MBA students enrolled in introductory corporate finance classes. They were told they would have an opportunity to earn money in a research experiment involving group decision making. Every subject participated in only one experimental session.

A. The Basic Design

Some basic features were common to all the treatments. Variations of the basic design are discussed later. First, subjects were read a set of instructions (see Appendix B); they completed assigned worksheets and were given the opportunity to ask questions. The terms “insiders,” “watchdogs,” and “project” were never mentioned in the instructions or used orally during the experiment. Instead, insiders were always referred to as “Type A” participants, watchdogs as “Type B” participants, and the project as the “action.” The term “participants” was used instead of “agents.”

At the end of the instructional period, the monitors randomly divided the subjects into groups of seven. Next, subjects were randomly assigned their agent type—insider or watchdog. To establish a benchmark against which the performance of watchdog-dominated boards could be assessed, in one treatment, all subjects were classified as insiders. The remaining treatments employed groups consisting of three insiders and four watchdogs. This insider–outsider split is in keeping with Jensen (1993), among others, who has suggested that board size should be limited to seven or eight members, so that the marginal cost of coordination and processing problems does not exceed the marginal benefit of the additional members’ input. In the context of our model, the group size and the number of insiders was the minimum number needed to ensure that (1) the defection of one insider from a unanimous vote by insiders did not cancel the majority insider vote and (2) outsiders had enough votes to override insiders.

Groups dispersed to different ends of a large classroom to commence the session. Each round began with a period of communication between group members. The following restrictions applied: No physical threats, no side payments, no communication among groups, and a maximum of four minutes for each discussion period. Subjects never appeared to find this time limit to be binding.

Next, the insiders from each group watched a monitor draw the project type from a bucket. To ensure that good and bad draws had equal probabilities, the bucket contained 50 white chips (good outcome) and 50 red chips (bad outcome) and chips were replaced after each draw. Following a draw, the insiders returned to their groups and a discussion among all members of each group commenced.

The discussions seldom approached the two-minute time limit. Private ballots were cast following this discussion. A monitor then counted the votes.

The project was undertaken if the yes votes outnumbered the no votes. When the yes votes equaled or were less than the no votes, the project was rejected. The monitor privately informed each group of the outcome and distribution of votes by agent types. Most treatments incorporated the following penalty feature for split votes: If at least one insider's vote did not conform with the majority vote for the group, a monitor drew a chip from a bucket of poker chips that contained 20 blue chips and 80 white ones. Chips were replaced after each draw. Insiders in the group were assessed a penalty if a blue chip was drawn. The outcome, together with the project type and the occurrence of a split vote, determined payoffs for the period.

In accordance with Assumption (3), payoffs were designed to ensure that insiders preferred to accept the project regardless of the outcome. Absent the penalty for lack of consensus, they received at least \$0.90 following project acceptance, compared with a maximum \$0.60 following its rejection. When a penalty was imposed, insider payoffs fell by \$0.25. Thus, insiders earned at least \$0.65 if the project was undertaken even if a penalty was imposed. Because this was higher than their expected \$0.60 payoff if the project was rejected and no penalty was imposed, in accordance with Assumption (5), the penalty was not sufficient to reverse their preferences between investing in the project and rejecting it.

Watchdogs' payoffs were designed to ensure that they preferred taking on the project only if it was good. They could expect to earn \$0.70 from project acceptance conditional on a good draw and \$0.00 from acceptance conditional on a bad draw. Consistent with Assumption (4), their expected payoff of \$0.35 from acceptance was less than their expected payoff of \$0.50 from rejection. All payoffs were common knowledge.

A round ended after subjects learned about the outcome, participated in the penalty draw if applicable, and calculated their earnings. Each experimental session consisted of 10 rounds, but subjects essentially played a game with an indefinite endpoint since they were not told how many rounds they would play.

B. The Central Treatments

We employed five central treatments to examine the importance of the two main features of the model described above—boards with watchdog majorities and penalties for split insider votes. As detailed in Table I, the treatments allowed for verbal communication between subjects both before and after the drawing for project quality. Pre-draw communication followed the group-subgroup sequence, that is, communication was permitted first among all members of a group, and then within subgroups based on agent type. After the drawing and before voting, “informed” insiders were allowed to communicate with “uninformed” watchdogs in their group.

In three treatments, we employed a random mixing protocol (RA) where, after each round, subjects were randomly assigned to new groups. However, they maintained their agent type for the entire session. Similar random mixing protocols

Table I
Description of the Central Treatments

This table describes the five central treatments, including the number of groups employing each treatment and the distribution of draws for each treatment. The treatments can be categorized based on the mixing protocols employed—random mixing (RA), where group membership was changed after every round but subjects retained their agent-type for the entire session, and repeated groups (RE), where group composition remained unchanged for the duration of the session. No watchdogs were included in treatment RANW. The remaining treatments included four watchdogs in each seven-member group. In treatments RANW, RAP, and REP, insiders faced the possibility of a penalty if even one of them voted against the majority. In treatments RANP and RENP, there was no penalty for split votes. The group-subgroup communication protocol (GS) was used in all treatments, that is, before the draw for project quality, all group members were allowed discussion time followed by discussion time within the subgroup of insiders and watchdogs.

	Treatment				
	RANW	RAP	RANP	REP	RENP
Mixing protocol	RA	RA	RA	RE	RE
Penalty for split votes	Yes	Yes	No	Yes	No
Number of watchdogs	0	4	4	4	4
Members in each group	7	7	7	7	7
Communication permitted	Yes	Yes	Yes	Yes	Yes
Mode of communication	Verbal	Verbal	Verbal	Verbal	Verbal
Communication protocol before project-quality draw	GS	GS	GS	GS	GS
Number of groups	3	3	2	3	2
Number of draws	30	30	20	30	20
Number of good draws	18	18	12	10	8
Average payoff to insiders (\$)	10.31	9.30	9.30	7.62	8.08
Average payoff to outsiders (\$)	NA	6.20	6.20	5.73	5.45

are employed extensively in the experimental literature to obtain multiple observations of a single shot game with a limited number of subjects (see, e.g., Forsythe et al. (1996)). The remaining two central treatments employed the repeated groups protocol (RE). Group membership and agent type remained stable for the duration of a session when this protocol was employed. This mixing protocol is designed to capture real world situations where board membership remains stable across a number of votes. This variation in mixing protocols was also motivated by extant research, which demonstrates that mixing protocols can influence experimental outcomes (see, e.g., Eckel and Holt (1989)).

In the first central treatment (RANW), we employed the random mixing protocol and no subjects were assigned the role of watchdogs, that is, each session involved seven insiders and no watchdogs. In all remaining treatments, four subjects were assigned the role of watchdogs and three to the role of insiders. The second central treatment, labeled RAP, employed the random mixing protocol and insiders were subject to a penalty for split insider votes. To examine the impact of the penalty, we ran sessions of the random mixing protocol where the insiders did not face a penalty draw after a split vote. This treatment is labeled RANP. Similarly, the repeated mixing protocol incorporating penalties for split

insider votes is labeled REP, while the repeated mixing protocol treatment without penalties for split insider votes is labeled RENP.

We planned to run three sessions for each treatment. However, even though we overrecruited by five subjects in the first nonpenalty experiment (RANP), not enough students showed up to form the third group of seven. Given our budget constraint and the fact we would have at least 20 observations with the two groups, we decided to go ahead and run the experiment. For consistency, we also ran only two sessions for treatment RENP.

C. Robustness

Experimental and theoretical research demonstrates that the structure of communication matters (see, e.g., Farrell and Rabin (1996) and Milgrom and Roberts (1996)). As Isaac and Walker (1988) suggest, communication may (1) speed agents' awareness of optimal group behavior—implying that the effects of communication will continue even if the ability to communicate is later removed, or (2) influence agents' beliefs about other agents' ongoing responses—implying that communication must continue if it is to be effective. Mirroring the recognition of the importance of communication in academic research, recent calls for reform of corporate boards have included suggestions aimed at changing the structure of communications among board members (see MacAvoy and Millstein (1999)).

Inspired by both existing research and suggestions for board reform, we ran nine treatments to examine the robustness of results from our central treatments. A summary of the features of the robustness treatments is presented in Table II. The first six treatments were identical to the central treatment REP in all respects except the communication protocols. In the first of these treatments, to provide a baseline for the effect of communication, no subject communication was permitted during the entire session. We ran three sessions of this treatment. Only one session was run for each of the remaining robustness treatments because these treatments were merely exploratory in nature. The second robustness treatment reversed the pre-draw sequence of communication in the central treatments, employing the subgroup-group protocol—subgroups of insiders and watchdogs first discussed strategies separately before joining their groups to continue the discussions.

Each of the next four robustness treatments varied the sequence of communication prior to the draw for project quality and did not permit verbal communication following the draw. In the first of these treatments, we allowed for an additional period of communication between agent subgroups before the project-quality draw—the subgroup-group-subgroup sequence of communication. The next robustness treatment employed the subgroup-group sequence communication protocol prior to the project-quality draw, just as did the second robustness treatment described above. The fifth treatment employed the group-subgroup sequence of communication employed in the central treatments. The sixth treatment only allowed for communication within the group as a whole, but prohibited subgroup communication.

Table II
Description of the Robustness Treatments

This table describes the robustness treatments including the number of groups employing each treatment and the distribution of draws for each treatment. Across all treatments, the seven-member groups included four watchdogs, and insiders faced the possibility of a penalty if even one of them voted against the majority. Only the mixing and communication protocols varied. The first six treatments employed the repeated groups (RE) mixing protocol, where group composition and agent types remained unchanged for the duration of the session. Treatments 7 through 9 employed the random watchdogs (RAW) protocol, where watchdogs were randomly shuffled across groups after every round but insiders remained with their groups. No communication was permitted in treatment 1. In the remaining treatments, subjects could communicate verbally prior to the project-quality draw. However, the sequence of discussion between the entire group (G) and subgroups of agents (S) prior to the project-quality draw varied across treatments. For example, treatment 3 employed the SGS protocol, that is, first subjects were permitted to communicate only with other agents of their type (watchdogs or insiders) after which the entire group was permitted to communicate, the communication period ended with communication once again being restricted to subgroups. Following the project-quality draw, in treatments 2 through 9, communication was permitted only during the first five rounds. In treatment 2, the post project-quality draw communication was verbal. In treatments 3 through 9, communication after the draw for project quality took the form of insiders passing watchdogs a piece of paper after having circled either “yes” or “no” on it.

	Treatment								
	1	2	3	4	5	6	7	8	9
Mixing protocol	RE	RE	RE	RE	RE	RE	RAW	RAW	RAW
Penalty for split votes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes	Yes
Number of watchdogs	4	4	4	4	4	4	4	4	4
Members in each group	7	7	7	7	7	7	7	7	7
Periods with communication after project quality draw	0	10	First 5	First 5	First 5	First 5	First 5	First 5	First 5
Mode of communication following project-quality draw	NA	Verbal	Message	Message	Message	Message	Message	Message	Message
Communication protocol before project-quality draw	NA	SG	SGS	SG	GS	G	GS	SG	G
Number of groups	3	1	1	1	1	1	1	1	1
Number of draws	30	10	10	10	10	10	10	10	10
Number of good draws	16	5	6	6	6	6	6	6	6
Average payoff to insiders (\$)	7.63	8.75	8.20	8.75	8.50	9.30	8.75	9.30	8.20
Average payoff to outsiders (\$)	4.93	6.00	5.80	6.00	6.00	6.20	6.04	6.04	6.04

In robustness treatments 3 through 6, for the first five rounds, following the project quality draw, communication took the form of the insider subgroup passing a piece of paper to the watchdog subgroup via a monitor. On the piece of paper, insiders could convey their strategy only by circling either yes or no. During the final five periods, no communication was permitted following the project-quality draw. Because a written message is not as easily misinterpreted, it should provide an environment more conducive to truth telling and reputation building than face-to-face communication. Thus, we analyzed the communication-mode protocol of a written message followed by rounds with no communication.

Robustness treatments 7 through 9, the final three robustness treatments, mirrored the robustness treatments 4, 5, and 6 described above with one exception—they employed a hybrid of the mixing protocols. After each round, watchdogs were randomly shuffled across groups while insider subgroups remained unchanged. This randomized-watchdog mixing protocol was employed to examine the performance of the boards when insiders have longer tenures than watchdogs, as can be expected of boards of corporations whose management is entrenched (see, e.g., Shleifer and Vishny (1989)).

III. Results from Central Treatments

We now examine results from the central treatments along four dimensions: (1) the incidence of the institutionally preferred outcome, (2) insider voting patterns, (3) watchdog voting patterns, and (4) the predictive success of the two competing equilibria. The results indicate that efficient equilibria enjoy greater predictive success, and the introduction of watchdog majorities greatly increased the incidence of the institutionally preferred outcome. Although the penalty for split votes does not appear to have a perceptible effect on insider voting, the repeated mixing protocol appears to encourage insiders to behave opportunistically. Occasionally, this behavior leads to voting patterns consistent with coalition-proof equilibria.

A. Hypotheses

To help assess the results of the experiments, we now describe the predictions of our model for parameter values that are equal to those employed in our experiments. When a board consists of four watchdogs and three insiders, Theorem 1 asserts that, in all coalition-proof equilibrium outcomes,

- a. All three insiders vote to reject the project under both information signals,
- b. At least three watchdogs vote to reject the project under both information signals, and
- c. The project is rejected with probability 1, and insider consensus is never violated.

Theorem 2, when translated into the specific context of our experiments, argues that efficient equilibrium outcomes all share the following characteristics.

- a. All three insiders vote for the project under information signal G and against the project under information signal B ,
- b. Watchdogs tend to be passive in that at least two watchdogs abstain when the information signal is G , and
- c. The project is accepted if and only if the information signal is G .

These definitions also classify vote vectors in treatments without penalties for split votes. However, for these treatments, additional vote vectors may support efficient and coalition-proof equilibria as the elimination of the penalty removes the requirement for consensus among insiders. For example, certain coalition-proof equilibria allow for at least three watchdogs to reject regardless of the draw while three insiders vote to accept.

Although we do not explicitly model a board with no watchdogs, it can readily be established that in the absence of watchdogs, the coalition-proof equilibria require that all insiders vote unanimously to accept the project regardless of the draw. This follows because the coalition of all insiders maximizes its payoff by accepting the project regardless of the signal. In efficient equilibria it is still the case that insiders unanimously vote to accept following a good draw and reject following a bad one.

B. Incidence of the Institutionally Preferred Outcome

Table III presents the frequency with which the institutionally preferred outcome occurred. In the treatment with no watchdogs (RANW), following a bad draw, the institutionally preferred outcome never occurred. In fact, the project was rejected only once in 30 draws. Surprisingly, this rejection followed a good draw. The presence of watchdog majorities increased the adoption of the institutionally preferred policy. In both treatments employing random mixing, the institutionally preferred outcome always prevailed. Its incidence fell off in the repeated group treatments, with the greatest decline following good draws. In treatment REP, following good draws, the institutionally preferred policy was adopted only 60 percent of the time compared with 95 percent following bad draws. This difference can be attributed primarily to one session of REP where the project was rejected three times after good draws. The drop in the acceptance rate following good draws, together with the relatively low acceptance rate following bad draws, supports the coalition-proof equilibria. However, the efficient equilibria continue to describe outcomes well, as the percentage of institutionally preferred outcomes in two of three sessions of treatment REP is comparable to that in the random mixing sessions.

We used Chi-squared statistics to examine differences in the outcome distributions among treatments. The results, presented in Table IV, confirm that watchdog majorities significantly influenced the adoption of the institutionally preferred policy following bad draws, suggesting that outsider-dominated boards

Table III**Adoption of the Institutionally Preferred Outcome in the Central Treatments**

This table presents the frequency with which the institutionally preferred outcome—accept if the project is good and reject if it is bad—was adopted in the central treatments. This information is presented for all project-quality draws and by type of draw. The number of draws of each type is presented in parentheses. Panel A presents this information for the random mixing treatments, while Panel B presents this information for the repeated groups treatments.

	Draw		
	All	Good	Bad
Panel A: Random Mixing Treatments			
No watchdogs (RANW)	17 out of 30 (56.7%)	17 out of 18 (94.4%)	0 out of 12 (0.0%)
Watchdogs and penalty (RAP)	30 out of 30 (100.0%)	18 out of 18 (100.0%)	12 out of 12 (100.0%)
Watchdogs no penalty (RANP)	20 out of 20 (100.0%)	12 out of 12 (100.0%)	8 out of 8 (100.0%)
Panel B: Repeated Groups Treatments			
Watchdogs and penalty (REP)	25 out of 30 (83.3%)	6 out of 10 (60.0%)	19 out of 20 (95.0%)
Watchdogs no penalty (RENP)	18 out of 20 (90.0%)	7 out of 8 (87.5%)	11 out of 12 (91.7%)

Table IV

Tests for Differences in the Incidence of the Institutionally Preferred Outcome across the Central Treatments

This table presents Chi-squared statistics for differences in the incidence of the institutionally preferred outcome—accept if the project is good and reject if it is bad—across the central treatments. For each test, outcomes are categorized as either institutionally preferred or not. Panel A provides Chi-squared statistics following good draws and Panel B provides statistics following bad draws. The term UD signifies that the Chi-squared statistic is undefined. The significance levels for the Chi-squared statistics are as follows: $\chi^2_{(1, 0.10)} = 2.71$, $\chi^2_{(1, 0.05)} = 3.84$, and $\chi^2_{(1, 0.01)} = 6.63$. Significance at the 5 percent and 1 percent confidence levels is denoted by ** and *, respectively.

	Treatment			
	Random No Watchdogs (RANW)	Random Penalty (RAP)	Random No Penalty (RANP)	Repeated Penalty (REP)
Panel A: Following a Good Draw				
Random penalty (RAP)	1.02			
Random no penalty (RANP)	0.70	UD		
Repeated penalty (REP)	5.17**	8.39*	5.86**	
Repeated no penalty (RENP)	0.38	2.37	1.58	1.67
Panel B: Following a Bad Draw				
Random penalty (RAP)	24.00*			
Random no penalty (RANP)	20.00*	UD		
Repeated penalty (REP)	28.08*	0.62	0.41	
Repeated no penalty (RENP)	20.32*	1.04	0.70	0.14

can significantly improve resource allocation. Contrary to expectations, however, the penalty for split votes had little effect. For both mixing protocols, there is no evidence that the elimination of the penalty significantly affected the incidence of institutionally preferred outcomes. The tests do, however, indicate significant differences across mixing treatments, suggesting that repeated interaction between agents may adversely affect board performance.

C. Insider Voting Patterns

Table V presents insider votes. A cursory examination reveals near universal support for the project in treatment RANW, strongly supporting the prediction that boards with no watchdogs are likely to misallocate resources. Only twice did an insider vote to reject. In each case the vote was cast after a good draw. The first no vote occurred after six rounds, when one subject convinced his group to vote against the project despite a good draw. This resulted in a unanimous vote to reject the project. A monitor overheard the subject suggest that there might be an additional reward for rejecting the project, even though the monitor had assured the group otherwise during the instruction period. In the next period the same individual voted no again, despite a good draw. The remainder of the group, however, voted to accept.

Table V
Insider Voting Patterns in the Central Treatments

This table presents insider voting patterns in the central treatments. Each cell presents the frequency of a combination of yes and no votes. Panel A presents voting patterns following good draws and Panel B presents voting patterns following bad draws. Column 5 in Panel A presents the total number of good draws and column 5 in Panel B presents the total number of bad draws.

	Voting Pattern				
	All	One	Two	All	Total
	Yes	No	No	No	Draws
	(1)	(2)	(3)	(4)	(5)
Panel A: Following a Good Draw					
Random mixing					
No watchdogs (RANW)	16	1	0	1	18
Watchdogs and penalty (RAP)	18	0	0	0	18
Watchdogs and no penalty (RANP)	12	0	0	0	12
Repeated groups					
Watchdogs and penalty (REP)	7	0	1	2	10
Watchdogs and no penalty (RENP)	8	0	0	0	8
Panel B: Following a Bad Draw					
Random mixing					
No watchdogs (RANW)	12	0	0	0	12
Watchdogs and penalty (RAP)	0	0	0	12	12
Watchdogs and no penalty (RANP)	0	0	0	8	8
Repeated groups					
Watchdogs and penalty (REP)	3	1	0	16	20
Watchdogs and no penalty (RENP)	5	0	0	7	12

The introduction of watchdog majorities altered insider voting. In the two random mixing treatments, insider behavior conformed perfectly with efficient equilibria, as insiders always voted unanimously to accept following a good draw and to reject following a bad one. The penalty for a lack of consensus did not alter voting.

Behavior in the repeated group treatments fell between the two extremes. The vast majority of insiders displayed some form of opportunistic behavior. In treatment REP, following bad draws, the majority of insiders voted to accept 20 percent of the time. In treatment RENP, following bad draws, insiders unanimously voted to accept 42 percent of the time. Further, across all sessions employing the repeated groups protocol, after bad draws, at least one insider voted to accept 28 percent of the time. Following good draws, in treatment REP, at least one insider voted to reject 30 percent of the time. In four of five sessions, a majority of insiders voted to accept the project at the first bad draw. In two sessions of treatment REP, an insider majority voted to reject following a good draw and to accept following a bad one, providing some support for coalition-proof equilibria. With the exception of one session each of treatments REP and RENP, however,

Table VI
Tests for Differences in Insider Voting Patterns across the Central Treatments

This table presents Chi-squared statistics for difference in the distribution of the insider votes across treatments. For each test, insider votes are classified into four groups based on the percentage of insider yes votes. The ranges for these four categories are as follows: 0 to 25 percent, 26 to 50 percent, 51 to 75 percent, and 76 to 100 percent. Panel A provides Chi-squared statistics following good draws, and Panel B provides this information following bad draws. The term UD signifies that the Chi-squared statistic is undefined. The symbol †† (†) denotes that two categories (three categories) of voting patterns were combined due to zero observations, resulting in only 2 (1) degrees of freedom. The significance levels for the Chi-squared statistics are as follows: $\chi^2_{(1, 0.10)} = 2.71$, $\chi^2_{(1, 0.05)} = 3.84$, $\chi^2_{(1, 0.01)} = 6.63$, $\chi^2_{(2, 0.10)} = 4.61$, $\chi^2_{(2, 0.05)} = 5.99$, $\chi^2_{(2, 0.01)} = 9.21$, $\chi^2_{(3, 0.10)} = 6.25$, $\chi^2_{(3, 0.05)} = 7.81$, and $\chi^2_{(3, 0.01)} = 11.34$. Significance at the 5 percent and 1 percent confidence levels is denoted by ** and *, respectively.

	Treatment			
	Random No Watchdogs (RANW)	Random Penalty (RAP)	Random No Penalty (RANP)	Repeated Penalty (REP)
Panel A: Following a Good Draw				
Random penalty (RAP)	2.12 ^{††}			
Random no penalty (RANP)	1.43 ^{††}	UD		
Repeated penalty (REP)	3.89	6.05 ^{††**}	2.63 ^{††}	
Repeated no penalty (RENP)	0.97 ^{††}	UD	UD	2.88 ^{††}
Panel B: Following a Bad Draw				
Random penalty (RAP)	24.00 ^{†*}			
Random no penalty (RANP)	20.00 ^{†*}	UD		
Repeated penalty (REP)	21.76 ^{††*}	2.74 ^{††}	1.87 ^{††}	
Repeated no penalty (RENP)	9.88 ^{†*}	6.34 ^{†**}	4.44 ^{†**}	3.22 ^{††}

following the first two draws, insider behavior resembled efficient equilibrium strategies.

Chi-squared test statistics to examine the effect of mixing protocols and penalties on insider voting are presented in Table VI. These tests highlight the impact watchdogs have on board performance, as they indicate that the presence of watchdog majorities significantly affected insider voting following bad draws. The tests also support the hypothesis that the mixing protocol can influence voting. Following bad draws, voting differed significantly between treatment RENP and treatments RAP and RANP. Following good draws, insiders votes were significantly different across treatments REP and RAP. As with the outcome distributions, we find no evidence that the penalty for split votes significantly altered the distribution of votes.

D. Watchdog Voting Patterns

We now turn to an examination of watchdog voting. Table VII shows that watchdogs usually voted unanimously to reject the project following bad draws and abstain following good draws, suggesting that insiders accurately transmitted information about project quality. Note also that when all watchdogs vote unanimously, no individual watchdog is pivotal in the vote, leaving all watchdogs indifferent between abstaining and voting against the project. This type of

Table VII
Watchdog Voting Patterns in the Central Treatments

This table presents watchdog vote distributions in the four central treatments that included watchdogs. Each cell presents the frequency of a combination of abstain and no votes. Panel A presents the frequency of votes following good draws and Panel B presents the frequency of votes following bad draws. Column 6 in Panel A presents the total number of good draws, while column 6 in Panel B presents the total number of bad draws.

	Voting Pattern					
	All Abstain	One No	Two No	Three No	All No	Total Draws
	(1)	(2)	(3)	(4)	(5)	(6)
Panel A: Following a Good Draw						
Random mixing						
Watchdogs and penalty (RAP)	18	0	0	0	0	18
Watchdogs and no penalty (RANP)	12	0	0	0	0	12
Repeated groups						
Watchdogs and penalty (REP)	4	3	0	1	2	10
Watchdogs and no penalty (RENP)	7	0	0	0	1	8
Panel B: Following a Bad Draw						
Random mixing						
Watchdogs and penalty (RAP)	0	0	1	0	11	12
Watchdogs and no penalty (RANP)	0	1	1	0	6	8
Repeated groups						
Watchdogs and penalty (REP)	5	1	2	0	12	20
Watchdogs and no penalty (RENP)	1	0	0	1	10	12

Table VIII
Tests for Differences in Watchdog Voting Patterns across the Central Treatments

This table presents Chi-squared statistics for differences in watchdog votes across the four central treatments that included watchdogs. For each test, watchdog votes are classified into five groups based on the number of watchdog no votes. Panel A provides Chi-squared statistics following good draws and Panel B provides statistics following bad draws. The term UD signifies that the Chi-squared statistic is undefined. In some instances, vote categories were combined because of zero observations, reducing the degrees of freedom for the tests. In some instances voting categories were combined because they contained no observations, reducing the number of degrees of freedom. The symbols †††, ††, and †, denote 3, 2, and 1 degrees of freedom, respectively. The significance levels for the Chi-squared statistics are as follows: $\chi^2_{(1, 0.10)} = 2.71$, $\chi^2_{(1, 0.05)} = 3.84$, and $\chi^2_{(1, 0.01)} = 6.63$; $\chi^2_{(2, 0.10)} = 4.61$, $\chi^2_{(2, 0.05)} = 5.99$, and $\chi^2_{(2, 0.01)} = 9.21$; $\chi^2_{(3, 0.10)} = 6.25$, $\chi^2_{(3, 0.05)} = 7.81$, and $\chi^2_{(3, 0.01)} = 11.34$; $\chi^2_{(4, 0.10)} = 7.78$, $\chi^2_{(4, 0.05)} = 9.49$, and $\chi^2_{(4, 0.01)} = 13.28$. Significance at the 5 percent and 1 percent confidence levels is denoted by ** and *, respectively.

	Treatment		
	Random Penalty (RAP)	Random No Penalty (RANP)	Repeated Penalty (REP)
Panel A: Following a Good Draw			
Random no penalty (RANP)	UD		
Repeated penalty (REP)	13.76 ^{†††*}	9.92 ^{†††**}	
Repeated no penalty (RENP)	2.36 [†]	1.58 [†]	4.99 ^{†††}
Panel B: Following a Bad Draw			
Random no penalty (RANP)	1.74 ^{††}		
Repeated penalty (REP)	3.98 ^{†††}	2.68 ^{†††}	
Repeated no penalty (RENP)	3.05 ^{†††}	4.38	5.17

indifference is common in voting games. Nevertheless, one might expect that, given such indifference, watchdog strategies may wander from the equilibrium. However, such a tendency was not observed.

Once again the mixing protocol appears to have influenced voting. In random mixing treatments, watchdog votes strongly supported the hypothesis that outsider-dominated boards implement the institutionally preferred outcome. In the repeated group treatments, watchdog votes were sometimes consistent with the coalition-proof equilibria.

Voting patterns varied following good draws. In treatments RAP and RANP, watchdogs always voted unanimously to abstain following good draws. In the repeated group treatments, however, watchdogs occasionally voted unanimously to reject the project following good draws, providing some support for the coalition-proof equilibria. In treatment RENP, unanimous rejection followed good draws 12.5 percent of the time. In treatment REP, following good draws, watchdogs unanimously rejected the project 20 percent of the time. Unanimous watchdog rejection followed soon after votes in which insiders managed to convince watchdogs to accept the project despite a bad draw.

Chi-squared tests presented in Table VIII indicate that changes in the mixing protocol influenced watchdog behavior. For treatments RAP and REP the difference in

behavior is significant following good draws. Once again, the tests do not support the hypothesis that the penalty for split votes influenced watchdog behavior.

E. Subject Votes and Equilibrium Vote Vectors

We now examine the consistency of aggregate vote vectors with each of the two sets of equilibria identified above. Table IX presents information on the consistency of voting patterns with equilibrium predictions. The criteria for categorizing the votes are described above. The use of a broader definition for equilibrium vote vectors in the treatments without penalties only changes the results for treatment RENP following bad draws. In this case, the proportion of vote vectors supporting the efficient (coalition-proof) equilibria rises from 58.3 percent (58.3 percent) to 91.7 percent (91.7 percent).

Across all treatments, subject votes were consistent with efficient equilibria 78 percent of the time. In the treatment without watchdogs (RANW), votes were never consistent with efficient equilibria following bad draws. In treatments with

Table IX
Consistency with Equilibrium Outcomes in the Central Treatments

This table presents the percentage of vote vectors that are consistent with each of two sets of equilibria—efficient equilibria and coalition-proof equilibria. The second column presents the percentage of votes that are consistent with efficient equilibria and the third column presents the percentage of votes that are consistent with coalition-proof equilibria. Panel A presents this data following good draws and Panel B presents this data following bad draws. Votes were classified as being consistent with efficient equilibria if insiders unanimously supported (rejected) the project following a good (bad) draw and, in treatments including watchdogs, at least two watchdogs abstained following good draws. Votes were classified as being consistent with coalition-proof equilibria if all insiders and at least three watchdogs voted to reject at all times. In the no watchdog treatment (RANW), votes were consistent with coalition-proof equilibria if all insiders voted to accept the project regardless of the draw.

	Equilibria	
	Efficient	Coalition-proof
Panel A: Following a Good Draw		
Random mixing		
No watchdogs (RANW)	16 out of 18 (88.9%)	16 out of 18 (88.9%)
Watchdogs and penalty (RAP)	18 out of 18 (100.0%)	0 out of 18 (0.0%)
Watchdogs and no penalty (RANP)	12 out of 12 (100.0%)	0 out of 12 (0.0%)
Repeated groups		
Watchdogs and penalty (REP)	6 out of 10 (60.0%)	2 out of 10 (20.0%)
Watchdogs and no penalty (RENP)	7 out of 8 (87.5%)	0 out of 8 (0.0%)
Panel B: Following a Bad Draw		
Random mixing		
No watchdogs (RANW)	0 out of 12 (0.0%)	12 out of 12 (100.0%)
Watchdogs and penalty (RAP)	12 out of 12 (100.0%)	11 out of 12 (91.7%)
Watchdogs and no penalty (RANP)	8 out of 8 (100.0%)	6 out of 8 (75.0%)
Repeated groups		
Watchdogs and penalty (REP)	16 out of 20 (80.0%)	9 out of 20 (45.0%)
Watchdogs and no penalty (RENP)	7 out of 12 (58.3%)	7 out of 12 (58.3%)

Table X
Chi-squared Statistics for Consistency of Votes with Equilibrium Outcomes in the Central Treatments

This table presents Chi-squared statistics for differences in the degree of consistency with equilibrium outcomes across the central treatments. Panel A presents statistics for insider vote distributions following good draws, Panel B presents statistics for watchdog vote distributions following good draws, and Panel C presents statistics for the aggregate vote (both insider and watchdog) vector following good draws. For each test, votes are classified into three groups—consistent with efficient equilibria, consistent with coalition-proof equilibria, and neither. Votes were classified as being consistent with efficient equilibria if insiders unanimously supported the project, and in treatments including watchdogs, at least two watchdogs abstained. Votes were classified as being consistent with coalition-proof equilibria if all insiders and at least three watchdogs voted to reject. The term UD signifies that the Chi-squared statistic is undefined. In some instances, categories were combined because zero observations reduced the degrees of freedom for the tests. The symbol †† denotes 2 degrees of freedom. The significance levels for the Chi-squared statistics are as follows: $\chi^2_{(1, 0.10)} = 2.71$, $\chi^2_{(1, 0.05)} = 3.84$, $\chi^2_{(1, 0.01)} = 6.63$, $\chi^2_{(2, 0.10)} = 4.61$, $\chi^2_{(2, 0.05)} = 5.99$, and $\chi^2_{(2, 0.01)} = 9.21$. Significance at the 10 percent and 5 percent confidence levels is denoted by *** and **, respectively.

	Treatment		
	Random Penalty (RAP)	Random No Penalty (RANP)	Repeated Penalty (REP)
Panel A: Insider Votes			
Random no penalty (RANP)	UD		
Repeated penalty (REP)	6.03 ^{††**}	4.17 ^{††}	
Repeated no penalty (RENp)	UD	UD	2.88 ^{††}
Panel B: Watchdog Votes			
Random no penalty (RANP)	UD		
Repeated penalty (REP)	6.05 ^{**}	4.17 ^{**}	
Repeated no penalty (RENp)	2.34	1.58	0.79
Panel C: Aggregate Vote Vectors			
Random no penalty (RANP)	UD		
Repeated penalty (REP)	8.46 ^{††**}	5.87 ^{††***}	
Repeated no penalty (RENp)	2.34	1.58	2.21 [†]

watchdogs, voting outcomes were very different. In treatments RAP and RANP, subject votes were consistent with the efficient equilibria 100 percent of the time. Confirming the earlier evidence on the impact of mixing protocols, voting behavior in the repeated group treatments was less consistent with efficient equilibria.

Following good draws, efficient equilibria were more successful in predicting subject strategies. Coalition-proof equilibria enjoyed some measure of success following bad draws, and appear to have no predictive power following good draws in treatments RAP and RANP. A simple explanation for this phenomenon is that, following bad draws, some of the efficient equilibria are supported by vote vectors that are identical to those supporting coalition-proof equilibria. Thus, following bad draws it is not always possible to allocate subject strategies to only one of the two competing equilibria.

Table X presents Chi-squared tests for differences in the success of competing equilibria across treatments. All vote vectors following good draws were classified as either consistent with efficient equilibria, coalition-proof equilibria, or neither. The tests for differences in these distributions across treatments confirm that voting in treatment REP is different from voting in treatments RAP and RANP. This supports our earlier results that the mixing protocol affects subject behavior. There is no evidence that the existence of penalties for a split vote affects the success of competing equilibria.

IV. Robustness

In this section we present evidence from our robustness sessions, focusing on the incidence of institutionally preferred outcomes and the consistency of outcomes with equilibria. The evidence indicates that our results are relatively insensitive to changes in communication protocols.

The frequency with which the institutionally preferred policy is adopted in the robustness treatments suggests that changes in communication protocols had

Table XI
Incidence of the Institutionally Preferred Outcome in the Robustness Treatments

This table presents the percentage of votes resulting in the institutionally preferred outcome—accept if the project is good and reject if it is bad—for one central treatment and nine robustness treatments. The central treatment presented is the repeated groups with penalty treatment (REP). The robustness treatments are classified first by the mode of communication following the project quality draw, then by the protocol governing the sequence of communication before the project quality draw. With the exception of the last three treatments in the table, they all employed the repeated groups mixing protocol. The final three treatments employed the random watchdogs mixing protocol.

	Draw		
	All	Good	Bad
Panel A: Central Treatment			
Repeated penalty (REP)	25 out of 30 (83.3%)	6 out of 10 (60.0%)	19 out of 20 (95.0%)
Panel B: Robustness Treatments			
No communication	19 out of 30 (63.3%)	9 out of 16 (56.3%)	10 out of 14 (71.4%)
Face to face			
Subgroup/group	10 out of 10 (100.0%)	5 out of 5 (100.0%)	5 out of 5 (100.0%)
Message-no message			
Group/subgroup	9 out of 10 (90.0%)	5 out of 6 (83.3%)	4 out of 4 (100.0%)
Subgroup/group	9 out of 10 (90.0%)	5 out of 6 (83.3%)	4 out of 4 (100.0%)
Group only	10 out of 10 (100.0%)	6 out of 6 (100.0%)	4 out of 4 (100.0%)
Subgroup/group/subgroup	8 out of 10 (80.0%)	4 out of 6 (66.7%)	4 out of 4 (100.0%)
Message-no message and random watchdogs			
Group/subgroup	9 out of 10 (90.0%)	5 out of 6 (83.3%)	4 out of 4 (100.0%)
Subgroup/group	9 out of 10 (90.0%)	5 out of 6 (83.3%)	4 out of 4 (100.0%)
Group only	9 out of 10 (90.0%)	5 out of 6 (83.3%)	4 out of 4 (100.0%)

Table XII

**Sensitivity of the Incidence of the Institutionally Preferred Outcome to
Changes in Mixing and Communication Protocols**

The table presents Chi-squared statistics for differences between the incidence of the institutionally preferred outcome—accept if the project is good and reject if it is bad—in the central treatment employing the repeated groups with penalties (REP) protocol and the robustness treatments. Statistics for outcome distributions conditioned on good draws and bad draws are presented in columns 2 and 3, respectively. The robustness treatments are classified first by the form of communication following the project quality draw, then by the protocol governing the sequence of communication prior to the project quality draw. With the exception of the last three treatments in the table, they all employed the repeated groups mixing protocol. The final three treatments employed the random watchdogs mixing protocol. Note that $\chi^2_{(1, 0.10)} = 2.71$, $\chi^2_{(1, 0.05)} = 3.84$, and $\chi^2_{(1, 0.01)} = 6.63$. Significance at the 10 percent confidence level is denoted by ***.

	Draw	
	Good	Bad
No communication	0.04	3.65***
Face to face		
Subgroup/group	2.73***	0.26
Message-no message		
Group/subgroup	0.96	0.22
Subgroup/group	0.96	0.22
Group only	3.20***	0.22
Subgroup/group/subgroup	0.073	0.22
Message-no message and random watchdogs		
Group/subgroup	0.96	0.22
Subgroup/group	0.96	0.22
Group only	0.96	0.22

little effect. Table XI indicates that, with the exception of the treatment with no communication, the adoption of the institutionally preferred policy in the robustness treatments uniformly exceeds that obtained in the central treatment employing the repeated groups with penalties (REP). The distributions of outcomes conditioned on project draws also produced this result. Statistical tests for differences in the outcomes of the robustness treatments and the central treatment REP, presented in Table XII, provide little support for the hypothesis that the outcome distributions are sensitive to changes in treatments.

Table XIII presents the frequency with which subject vote vectors were consistent with equilibrium vote vectors. In the robustness treatment without subject communication, both sets of equilibria displayed diminished explanatory power. Votes were consistent with the efficient equilibria only 30 percent of the time. Subject behavior could not be explained by the coalition-proof equilibria. This supports the idea that in the absence of communication subjects more likely will gravitate to equilibria that require less coordination—the efficient equilibrium. It also highlights the role of communication, even cheap talk, in helping subjects coordinate to equilibrium strategies.

Table XIII

Consistency with Equilibrium Outcomes in the Robustness Treatments

This table presents the percentage of votes in the robustness treatments that are consistent with each of two sets of equilibria—efficient equilibria and coalition-proof equilibria. Each cell in the second column presents the percentage of votes that is consistent with efficient equilibria. The third column documents consistency with coalition-proof equilibria. Panel A presents data following good draws and Panel B presents data following bad draws. Efficient equilibria call for unanimous insider support following a good draw and rejection following a bad one. At least two watchdogs abstain following a good draw. In coalition-proof equilibria, all insiders and at least three watchdogs vote to reject at all times.

	Equilibria	
	Efficient	Coalition-proof
Panel A: Following Good Draws		
No communication	9 out of 16 (56.3%)	0 out of 16 (0.0%)
Face to face		
Subgroup/group	5 out of 5 (100.0%)	0 out of 5 (0.0%)
Message-no message		
Group/subgroup	5 out of 6 (83.3%)	0 out of 6 (0.0%)
Subgroup/group	5 out of 6 (83.3%)	1 out of 6 (16.7%)
Group only	6 out of 6 (100.0%)	0 out of 6 (0.0%)
Subgroup/group/subgroup	4 out of 6 (66.7%)	2 out of 6 (33.3%)
Message-no message and random watchdogs		
Group/subgroup	5 out of 6 (83.3%)	1 out of 6 (16.7%)
Subgroup/group	5 out of 6 (83.3%)	0 out of 6 (0.0%)
Group only	5 out of 6 (83.3%)	1 out of 6 (16.7%)
Panel B: Following Bad Draws		
No communication	1 out of 14 (7.1%)	0 out of 14 (0.0%)
Face to face		
Subgroup/group	5 out of 5 (100.0%)	5 out of 5 (100.0%)
Message-no message		
Group/subgroup	4 out of 4 (100.0%)	4 out of 4 (100.0%)
Subgroup/group	4 out of 4 (100.0%)	3 out of 4 (75.0%)
Group only	4 out of 4 (100.0%)	4 out of 4 (100.0%)
Subgroup/group/subgroup	4 out of 4 (100.0%)	3 out of 4 (75.0%)
Message-no message and random watchdogs		
Group/subgroup	3 of 4 (75.0%)	1 of 4 (25.0%)
Subgroup/group	4 of 4 (100.0%)	3 of 4 (75.0%)
Group only	4 of 4 (100.0%)	2 of 4 (50.0%)

In the remaining treatments, once again efficient equilibria are successful in explaining subject behavior. At least 67 percent of votes resembled those supporting efficient equilibria. Votes were consistent with coalition-proof equilibria about 30 percent of the time. In most cases these votes followed bad draws. Again, the apparent power of coalition-proof equilibria to explain votes following bad draws might result because in both coalition-proof equilibria and some efficient equilibria, all insiders and a majority of watchdogs vote to reject. However, as with the central treatments, voting patterns following good draws provide some support for coalition-proof equilibria.

V. Conclusion

We model a board where a group of insiders, whose incentives are not aligned with those of the institution, has private information regarding the institutionally preferred allocation. Our laboratory analysis finds that the inclusion of a majority of uninformed watchdogs reduces the incidence of undesirable equilibria. Inefficient equilibria, however, cannot be completely eliminated. Institutionally preferred allocations are more likely to arise when group membership is frequently altered. In addition, the evidence suggests that a board size of seven, as suggested by Jensen (1993), among others, is sufficient for efficient decision making.

Gilson and Kraakman (1991) suggest that independent directors, to be effective, should not be merely independent of management but accountable to shareholders. Empirical studies use proxies for the alignment of incentive structures of independent directors with the institution, which are, by their nature, noisy. For example, though some directors may be classified as independent of management, they may be beholden to management in subtle ways, such as by acting as paid advisors or consultants to the company. Our experiments contribute to the research in this area by providing a framework that can control for such conflicts of interest and ensure that the independent directors are, truly, watchdogs.

The role of watchdog agents in many organizational structures and in our study has been primarily to vote on designated issues, but the positive role watchdogs play in attaining the institutionally preferred allocation in this governance structure may extend to a broader agenda-setting context. Other experimental studies have examined specific agenda-setting issues. An interesting area of further study would be to extend the research in both these areas to the role of independent directors within various agenda-setting environments.

Appendix A

This appendix provides the formal proofs of Theorems 1 and 2. Before we initiate our proofs we require some additional notation and definitions.

An agent's voting strategy is a map from observed messages into votes. Let M^N represent a vector of messages sent by the agents, $\mathbf{v}_i^W$ a voting strategy for watchdogs, and $\mathbf{v}_i^I$ a voting strategy for insiders. A strategy for an individual agent is thus an ordered pair of communication and voting strategies. We represent a strategy of watchdog i , by $\sigma_i^W \equiv (\mathbf{m}_i^W, \mathbf{v}_i^W)$, and the strategy of insider i , by $\sigma_i^I \equiv (\mathbf{m}_i^I, \mathbf{v}_i^I)$. Finally, let σ^I represent the vector of insider strategies; let σ^W represent the vector of watchdog strategies.

Strategies yield votes via the following functional compositions. The message strategies of the insiders and watchdogs, and the information signal, $s \in \{G, B\}$, will produce a pattern of messages. These messages will, when composed with the voting strategies of the watchdogs, produce the watchdog vote vector; when composed with insider voting strategies, they produce the insider vote vector. This process of composition generates a map, from strategies to votes, which we

represent as follows:

$$V(\sigma^W, \sigma^I(s)) = (\mathbf{v}^W((\mathbf{m}^W, \mathbf{m}^I(s))), \mathbf{v}^I((\mathbf{m}^W, \mathbf{m}^I(s)), s)) \quad (\text{A1})$$

Using the definition of agent utility given in (1) and (2), we see that the utility of an agent j under strategy vector σ is given by

$$u_j(\sigma) = E(U_j(\mathbf{V}(\sigma^W, \sigma^I), \omega)) \quad (\text{A2})$$

A. Definitions

DEFINITION 2: A strategy σ'_j is a best response to $\sigma' \in \Sigma$ if the following implication holds for all $\sigma \in \Sigma$

$$\forall k \in [N] - \{j\}, \sigma_k = \sigma'_k \Rightarrow u_j(\sigma') \geq u_j(\sigma) \quad (\text{A3})$$

DEFINITION 3: A strategy vector σ^* is a Nash equilibrium if for all $j \in [N]$, σ_j^* is a best response for j to σ^*

DEFINITION 4: A coalition-proof Nash equilibrium is defined by induction on the size of coalitions as follows. $S \subset [N]$

- (i) Suppose $\#(S) = 1$, then $S = \{j\}$ for some $j \in [N]$. In this case, σ is optimal for $S = \{j\}$ if and only if σ_j is a best response to σ for j .
- (ii) Assume optimality has been defined for all S such that $\#(S) \leq k - 1$. Define optimality of coalitions S of size k as follows:
 - (a) σ is self-enforcing for S if σ is optimal for T , whenever T is a strict subset of S .
 - (b) σ is optimal for S if it is self-enforcing for S and there does not exist any strategy vector σ' which is also self-enforcing for S such that

$$\forall j \in [N] - S, \sigma'_j = \sigma_j \quad (\text{A4})$$

$$\forall j \in S, u_j(\sigma') > u_j(\sigma) \quad (\text{A5})$$

Finally, if σ is optimal for $[N]$, we say that σ is a coalition-proof Nash equilibrium.

DEFINITION 5: If σ is a strategy vector and K is a subset of agents containing at least one insider, then we say that the insiders in K are decisive for σ for signal s if, holding the actions of all other agents fixed, they can force acceptance of the

project. In other words,

$$\begin{aligned}
 & \#(K \cap [I]) \\
 & + \#\{j \in [I] - K : \mathbf{V}_j(\sigma^W, \sigma^I(s)) = \mathcal{Y}\} \\
 & > \#\{j \in [w] : \mathbf{V}_j(\sigma^W, \sigma^I(s)) = \mathcal{N}\} \\
 & + \#\{j \in [I] - K : \mathbf{V}_j(\sigma^W, \sigma^I(s)) = \mathcal{N}\}.
 \end{aligned} \tag{A6}$$

Insiders belonging to a subset of agents are decisive for a given information signal if, by collectively changing their vote to yes when they receive that signal, they can ensure that the project is accepted. Next note that all insiders have the same payoff function and that all watchdogs have the same payoff function. Thus, the strategy vector that maximizes the payoff to a subset of agents that consists only of insiders or outsiders is well defined. This motivates the following definition.

DEFINITION 6: Let σ be a strategy vector; let K be a nonempty “pure” subset of agents consisting only of insider types or only of outsider types. Suppose that overall strategy vectors σ' such that $\sigma'_j = \sigma_j$, $j \notin K$, σ produces the highest payoff to agents in K ; then we say that σ maximizes payoffs over K .

B. Lemma Used to Establish Theorems 1 and 2

Our most important lemma, Lemma 1, is quite straightforward. It implies that in coalition-proof outcomes, pure coalitions consisting of just insiders or just outsiders act as if they are a single agent, maximizing their collective payoff over their joint strategy space.

LEMMA 1: Let σ be a strategy vector and let K be a nonempty “pure” subset of agents consisting only of insiders types or only of outsider types. The strategy σ is optimal for K if and only if σ maximizes payoffs over K .

Proof: Our proof is based on induction on the size of the coalition. If the coalition size, which we represent by k equals one, the lemma follows from the definition of a Nash equilibrium.

Next, suppose that Lemma 1 holds for a subset of size less than or equal to k . Consider a pure subset K of size $k+1$ and a strategy vector, σ , that maximizes type payoffs over K . All subsets of a pure subset must be pure. Maximizing a type's payoff over a subset of K can never yield a higher payoff than the payoff from maximizing over K . Thus, σ must be self-enforcing for $k+1$. Because σ maximizes over K , no other strategy vector produces a higher payoff. Thus, σ is optimal for K .

To prove the other leg of the if-and-only-if assertion, suppose that K is a pure coalition of size $k+1$ and let σ be a strategy vector that does not maximize type

payoffs over K . Then there must exist a strategy vector, say $\sigma' \neq \sigma$, such that σ' equals σ for agents not in K and σ' maximizes the payoff over K . (Because there are only a finite number of distinct strategy vectors, existence of a maximizing vector is guaranteed.) By the results of the previous paragraph, σ' is optimal, and thus, *a fortiori*, self-enforcing for K . This implies that σ cannot be optimal for K . Thus, maximization of the type payoffs over strategies in K is a necessary condition for optimality as well. Q.E.D.

By Lemma 1, insiders will force project acceptance except when project acceptance is blocked by insufficient watchdog votes. Because information regarding the information signal is transmitted only by informed insiders, insiders can always force acceptance under one signal if they can force acceptance under any signal. This reasoning underlies the next lemma, Lemma 2.

LEMMA 2: *If σ is any coalition-proof Nash equilibrium, σ , the project is accepted with probability 1 or probability 0.*

Proof: Suppose that, under σ , the project is accepted under signal s_1 but not under signal s_2 . Consider the subset consisting of all insiders. Consider the strategy vector σ' calling for insiders to use the message and voting strategy that they use when receiving s_1 under σ for both information signals. Under σ' the project is accepted with probability 1. Assumptions (3), (4), and (5) ensure that this outcome produces a higher insider payoff than strategy σ . Thus, σ cannot maximize the payoffs over $[I]$. Thus, σ is not optimal for $[I]$ and thus is not coalition-proof. Q.E.D.

Because the payoff to watchdogs is always higher if the project is rejected under both signals than it is if the project is accepted under both signals, the coalition of all watchdogs can always gain by forcing universal rejection in any candidate equilibrium in which the project is being accepted with probability 1. Thus, such equilibria are not coalition-proof.

LEMMA 3: *In any coalition-proof equilibrium, insiders are not decisive under either information signal, that is, under both signals watchdogs cast sufficient votes against the project to block passage regardless of the votes of the insiders.*

Proof: Consider a coalition-proof Nash equilibrium σ . By Lemma 2 we know that if the project is accepted at all, then the project is accepted with probability 1. Thus, the equilibrium payoff to watchdogs is $\pi x(W, A, G) + (1 - \pi)x(W, A, B)$.

Now suppose watchdogs deviate to the strategy of always voting against the project, that is, consider the strategy vector σ' defined as follows. For insiders, play the strategies prescribed by σ ; for outsiders, play the message strategies prescribed by σ but follow the voting strategy of voting against the project regardless of the message sent in the message phase. Because watchdogs outnumber insiders, the project is rejected with probability 1. Thus, the strategy vector σ' yields watchdogs a payoff of $\pi x(W, R, G) + (1 - \pi)x(W, R, B)$. By (4), this exceeds the equilibrium payoff under σ . Thus, σ does not maximize watchdog payoffs

over $[w]$. By Lemma 1, σ is not optimal for $[w]$ and thus σ is not coalition-proof. Q.E.D.

Because the project is being rejected in all coalition-proof equilibria regardless of how insiders vote, and because of the penalty imposed on insiders when consensus fails, insiders collectively have an incentive to vote unanimously against project acceptance when their votes are not decisive. Thus, coalition-proof outcomes are characterized by unanimous insider rejection.

LEMMA 4: *In any coalition-proof equilibrium, all insiders vote to reject the project.*

Proof: To obtain a contradiction, let σ be a coalition-proof equilibrium in which not all insiders vote to reject the project. From Lemma 3 we know that in any coalition-proof equilibrium the project is voted down regardless of the insiders' voting behavior. Consider the subset of agents consisting of all insiders. If these insiders deviate to a strategy of voting against acceptance regardless of the messages sent in the message phase, the deviant strategy, by eliminating the possibility of the penalty for a lack of consensus, produces a higher payoff than the strategies insiders are playing under σ . Hence, σ does not maximize payoffs for $[I]$. Thus, by Lemma 1, σ is not optimal for $[I]$ and hence σ is not coalition proof. Q.E.D.

Lemmas 1 to 4 characterize coalition-proof equilibria. We now turn our attention to proving that coalition-proof equilibria exist. The existence proof requires us to consider mixed insider–watchdog coalitions. Characterizing such coalitions motivates the following definitions.

DEFINITION 7: *A coalition of agents, K , is flawed under strategy vector σ if there exists an information signal s such that the following conditions hold.*

a. *The project is rejected under s ; that is, for some s ,*

$$\#\mathcal{Y}(\mathbf{V}(\sigma^W, \sigma^I(s))) \leq \#\mathcal{N}(\mathbf{V}(\sigma^W, \sigma^I(s)))$$

b. *The coalition K contains at least one insider.*

c. *The insiders in K are decisive for σ under s .*

A flawed coalition contains a subset of insider agents who are decisive for project acceptance yet fail to ensure that the project is always accepted. In the subsequent analysis we will show that equilibria in which the set of all agents is flawed are not coalition proof. Because coalition proofness is defined by induction on subset size, we must define flawed coalitions not only for the set of all agents, but also for all proper subsets of agents. The next lemma shows that, when the strategy vector is flawed, by collective changes in their strategy, a sufficiently large coalition of insiders can always modify their strategies to ensure project acceptance.

LEMMA 5: *If J is flawed for σ and if all insiders not in J are sending the same message under both signals, there exists a strategy vector, σ' , which specifies the same strategies as σ for all watchdogs and insiders not in J such that the project is accepted with probability 1.*

Proof: We construct the strategy vector as follows. Let $K = [I] \cup J$, for all agents not in K ; let $\sigma' = \sigma$. Next, determine the signal, say s' , under which insiders are decisive. Each agent in K should (a) send under both signals the message that, under σ , she sent under s' and (b) subsequently vote to accept the project regardless of the pattern of messages received. This strategy will ensure that the project is accepted under both signals with the unanimous support of insiders in K . Q.E.D.

LEMMA 6: *If J is flawed for σ and if all insiders not in J are sending the same message under both signals, then σ is not optimal for J .*

Proof: By Lemma 5, σ does not maximize the payoff over $[I] \cup J$. Thus, by Lemma 1, σ is not optimal for J . Q.E.D.

LEMMA 7: *There exists a coalition-proof Nash equilibrium under which the project is rejected with probability 1 and all insiders cast votes against project acceptance under both information signals.*

Proof: Consider the strategy vector σ defined as follows. All agents sent the same arbitrary message, m_o , independent of the information signal, that is, $\mathbf{m}_i^W = \mathbf{m}_i^I(s) = m_o$. All insiders and watchdogs follow the strategy of voting $\mathcal{N}$ regardless of the information signal, that is, $\mathbf{v}_i^W(m) = \mathcal{N}$, $\mathbf{v}_i^I(m, G) = \mathcal{N}$, $\mathbf{v}_i^I(m, B) = \mathcal{N}$. To show that this is a coalition-proof Nash equilibrium, we need to show that σ is optimal for all subsets of $[N]$. To show this we provide a proof by induction on subset size. First note that when the subsets contain a single element, optimality simply requires that for all agents j , σ_j is a best response to σ for j . However, no individual agent can change the project acceptance decision by unilaterally changing his strategy. Moreover, insiders may call down a penalty if they deviate from the consensus. Thus, the assertion of optimality holds for all subsets of size one. Next suppose that, for subsets of size less than or equal to k , optimality holds. Consider a subset K of size $k + 1$. Given the induction hypotheses and the definition of coalition proofness, optimality requires that there does not exist another self-enforcing strategy for K that yields all the agents in K a higher payoff. If K consists only of watchdogs, K cannot produce a higher payoff because, given the uninformative messages of insiders, watchdogs cannot induce a strategy vector that accepts the project under the good information signal and rejects the project under the bad signal. Given assumptions (3), (4), and (5), rejecting the project under both signals produces a higher payoff to watchdogs than accepting the project under both signals. No improving vector of strategies exists for K , and thus, *a fortiori*, no self-enforcing vector of strategies exists for K when K consists of a set of watchdogs. Next note that a coalition of all insiders does not have sufficient votes

to change the outcome. Thus, such a subset cannot increase its welfare by deviating from the equilibrium. Only a mixed coalition, by implementing a vector of strategies calling for rejection when the information signal is B and acceptance when the information signal is G , can increase the payoff to all agents in K . However, by our earlier definition, such a coalition is flawed. Lemma 6 shows that a flawed coalition is not optimal and thus, *a fortiori*, is not self-enforcing. Thus, optimality for coalitions of size $k + 1$ has been established, proving the assertion of the theorem by induction. Q.E.D.

Proof of Theorem 1: This theorem follows directly from Lemmas 4, 5, 6, and 7. Q.E.D.

Proof of Theorem 2: First note that consensus among insiders is necessary in all Nash equilibria that implement the efficient outcome. This follows because a lack of unanimity among insiders would result in the possibility that insiders would bear a penalty. Thus, in all equilibria that implement the efficient outcome, all insiders vote $\mathcal{Y}$ if they observe the signal G , and vote $\mathcal{N}$ if they observe the signal B .

Next we show that an equilibrium exists that implements the efficient outcome. The equilibrium is given as follows. All agents send the same arbitrary message m_o , that is, $\mathbf{m}_i^W = \mathbf{m}_i^I(s) = m_o$. All insiders follow the strategy of voting $\mathcal{Y}$ if and only if they receive information signal G ; $\mathbf{v}_i^I(m, G) = \mathcal{Y}$, $\mathbf{v}_i^I(m, B) = \mathcal{N}$. All watchdogs abstain regardless of the messages they observe, $\mathbf{v}_i^W(m) = \mathcal{A}$. In this candidate equilibrium, the vote of an individual agent cannot change the project selected. Moreover, for insiders, deviation from the strategy may incur the penalty for lack of consensus. Thus, unilateral deviations from the candidate information strategy vector cannot increase the payoff to any of the agents. It follows that the candidate strategy vector is a Nash equilibrium. Q.E.D.

Appendix B

This appendix contains instructions for the *random-group* sessions with the penalty mechanism. The instructions for the other treatments are the same with obvious changes.

Instructions

GENERAL

You are about to participate in an experiment of group decision making. If you understand the instructions and make careful decisions, you may earn a considerable amount of money. These earnings will be paid to you, in cash, at the end of the experimental session.

There will be a number of decision-making periods. In each period, you will vote on whether or not to take on an **action**. A draw will be made, that together with the majority vote by members of your group, will determine the outcome for the period. Your earnings will be determined by the outcome, the occurrence of a split vote, and your participant type.

In each group, three members will be randomly chosen to be Type A participants, while the other four will be designated, Type B. Prior to voting in each period, there will be a discussion time among all seven members of your group. Next, the Type A participants and the Type B participants will have a discussion time in their respective subgroups. During these discussion periods all aspects of the voting choices may be discussed with the following restrictions:

- 1. there will be no discussion of physical threats;
- 2. there will be no discussion of side payments;
- 3. there will be no discussion between groups;
- 4. there will be a maximum of four minutes per discussion session.

Following the subgroup discussion time, the Type A participants in each group will observe the draw for that period. All members of your group will again meet together to discuss the choices. Voting will then take place **privately** by ballot slips which will be picked up by a monitor. The Type A participants can cast either a vote of “yes” or “no” for the action to be taken; while the Type B participants (who do not witness the draw) can cast either a vote to “abstain” from voting or place a vote of “no” for the action to be taken. After voting, the monitor will return to each group their group’s majority vote and a breakdown of votes by participant type. Earnings will be calculated and you will then be randomly selected into a different group for the next period. However, you will stay the same participant type throughout the experiment.

Draw:

Each period, the **payoff of the draw** will be determined by a participant randomly drawing a poker chip from a bucket. The bucket will contain 50 white chips and 50 red chips. A **white chip** represents **Draw I** while a **red chip** represents **Draw II**. After each drawing the chip will be returned to the bucket, thus Draw I and Draw II have an **equal** chance of being selected in each period. [Stop for demonstration of drawing chips from the bucket.]

Majority vote:

Whether the action is taken on or not for any group depends on the **majority vote** of *that* group.

A **majority vote** consists of more of one type of vote over the other type. A vote of **abstain** is a neutral vote that allows other members of the group (who do not abstain from voting) to determine whether the action is taken on. For example, a majority vote of “yes” (for the action to be taken) **must** consist of one of the following combinations of votes within a group:

<u>Yes</u>		<u>No and/or Abstain</u>
Three	with	at least two of the Type B participants abstaining
Two	with	all four of the Type B participants abstaining

Any other combination of votes will result in a no majority vote. Also, note that the tie votes (three yes, three no, and one abstain) or (two yes, two no, and three abstain) imply the action will **not** be taken.

Penalty:
A penalty will only affect the Type A participants. If **at least one** or more of the Type A participants within a *given* group do not vote with the majority vote of their group, there is a **one-fifth** chance that **all** Type A participants in *that* group will receive a penalty. Whether the penalty occurs or not will be determined by a random drawing from a bucket containing 40 white chips and 10 blue chips. A **blue chip** means that the **penalty will be in effect** while a **white chip** means the **penalty will not occur**. [Stop for demonstration of drawing chips from the bucket.]

EARNINGS
Your earnings depend on whether you are a Type A or Type B participant. Below are the payoffs for each participant type for all possible outcomes and penalties. The Type B participants always face the same payoffs. Note that Type A participants have two cases of potential payoffs: Case 1—no penalty payoffs and Case 2—penalty payoffs (that are \$0.25 less than the payoffs in Case 1).

PAYOFFS

Type A Participant:					
CASE 1: No Penalty			CASE 2: Penalty of \$.25		
Draw			Draw		
	I	II		I	II
Majority Vote			Majority Vote		
YES:	\$ 1.15	\$.90	YES:	\$.90	\$.65
NO:	\$.60	\$.60	NO:	\$.35	\$.35
Type B Participant:					
Draw					
	I	II			
Majority Vote					
YES:	\$.70	\$ 0			
NO:	\$.50	\$.50			

WORKSHEET I

Fill in the blanks below. If you have any questions, please raise your hand and a monitor will come to assist you.

A. No Penalty for Type A participants—use Case 1 payoffs.

	Your Vote	Other Group Members' Vote	Majority Vote	Action Outcome	Type A Earnings	Type B Earnings
1)	Yes	2 Yes; 0 No; 4 Abstain	—	II	—	—

2)	Abstain	3 Yes; 2 No; 1 Abstain	—	I	—	—
3)	No	0 Yes; 2 No; 4 Abstain	—	I	—	—
B. At least one Type A participant does not vote with the majority, and a blue chip is drawn indicating the penalty occurs—use Case 2 payoffs						
4)	Yes	1 Yes; 1 No; 4 Abstain	—	I	—	—
5)	No	2 yes; 1 No; 3 Abstain	—	I	—	—
6)	Abstain	3 yes; 3 No; 0 Abstain	—	II	—	—

WORKSHEET II

- 1) If the Type A participants in a group vote according to the actual draw then to earn the highest payoff the Type B participants, as a subgroup, should all vote —
- 2) If the Type A participants in a group do **not** vote according to the actual draw then to avoid the lowest payoff the Type B participants, as a subgroup, should all vote —

RECORD-KEEPING PROCEDURES

If you will now look at your record sheet you will see the following entries: GROUP NUMBER, PARTICIPANT TYPE, YOUR VOTE, MAJORITY VOTE, and DRAW. Each decision-making period you will record this information on your record sheet and also on your ballot. After the majority vote and the outcome of the action have been revealed to each group then each participant will calculate YOUR EARNINGS and record their CUMULATIVE EARNINGS. A monitor will come by to check your calculations.

Note that, your earnings are not affected by the decisions in any other group.

Please keep accurate records throughout the experiment and do not show another participant your record sheet.

PARTICIPANT TYPE

We will now randomly draw for participant type. The monitor will come to each group with a bowl containing seven poker chips. There will be three blue and four red chips. Those who draw a **blue chip** will be **Type A** participants and those who draw a **red chip** will be **Type B** participants.

THIS IS THE END OF THE INSTRUCTIONS. IF YOU HAVE ANY QUESTIONS PLEASE RAISE YOUR HAND AND ASK THEM AT THIS TIME. Please do not talk unless it is a designated discussion time.

References

Baysinger, Barry, and Henry Butler, 1985, Corporate governance and the board of directors: Performance effects of changes in board composition, *Journal of Law, Economics, and Organization* 1, 101–124.

- Bernheim, Douglas B., Bezalel Peleg, and Michael D. Whinston, 1987, Coalition-proof Nash equilibria I: Concepts, *Journal of Economic Theory* 42, 1–12.
- Bhagat, Sanjai, and Bernard S. Black, 1999, The uncertain relationship between board composition and firm performance, *The Business Lawyer* 54, 921–963.
- Byrd, John, and Kent Hickman, 1992, Do outside directors monitor managers: Evidence from tender offer bids, *Journal of Financial Economics* 32, 195–221.
- Cooper, Russell W., Douglas V. DeJong, Robert Forsythe, and Thomas W. Ross, 1989, Communication in the battle of the sexes game: Some experimental results, *RAND Journal of Economics* 28, 569–587.
- Eckel, Catherine, and Charles Holt, 1989, Strategic voting and agenda-controlled committee experiments, *American Economic Review* 79, 763–773.
- Farrell, Joseph, and Matthew Rabin, 1996, Cheap talk, *Journal of Economic Perspectives* 10, 103–118.
- Forsythe, Robert, Russell Lundholm, and Thomas Rietz, 1999, Cheap talk, fraud, and adverse selection in financial markets: Some experimental evidence, *Review of Financial Studies* 12, 481–518.
- Forsythe, Robert, Thomas Rietz, Roger Myerson, and Robert Weber, 1996, An experimental study of voting rules and polls in three-candidate elections, *International Journal of Game Theory* 25, 355–383.
- Gilson, Ronald, and Reinier Kraakman, 1991, Reinventing the outside director: An agenda for institutional investors, *Stanford Law Review* 43, 863–906.
- Hermalin, Benjamin, and Michael Weisbach, 1991, The effects of board composition and direct incentives on firm performance, *Financial Management* 20, 101–112.
- Isaac, Mark, and James M. Walker, 1988, Communication and free-riding behavior: The voluntary contribution mechanism, *Economic Inquiry* 26, 585–608.
- Jensen, Michael, 1993, The modern industrial revolution, exit, and the failure of internal control systems, *Journal of Finance* 48, 831–880.
- MacAvoy, Paul W., and Ira M. Millstein, 1999, The active board of directors and its effect on the performance of the large publicly traded corporation, *Journal of Applied Corporate Finance* 11, 8–20.
- Mehran, Hamid, 1995, Executive compensation structure, ownership and firm performance, *Journal of Financial Economics* 38, 163–184.
- Mikkelsen, Wayne H., and Megan Partch, 1997, The decline of takeovers and disciplinary managerial turnover, *Journal of Financial Economics* 44, 205–228.
- Milgrom, Paul, and John Roberts, 1996, Coalition-proofness and correlation with arbitrary communication possibilities, *Games and Economic Behavior* 17, 113–128.
- Palfrey, Thomas R., 1992, Implementation in Bayesian equilibrium, in J. -J. Laffont, ed.: *Advances in Economic Theory* (Cambridge University Press, Cambridge).
- Palfrey, Thomas R., and Sanjay Srivastava, 1993, *Bayesian Implementation* (Harwood Academic Publishers, Chur, Switzerland).
- Rietz, Thomas, Roger Myerson, and Robert Weber, 1998, Campaign finance levels as coordinating signals in three-way experimental elections, *Economics and Politics* 10, 185–217.
- Sefton, Martin, and Abdullah Yavas, 1996, Abreu-Matsushima mechanisms: Experimental evidence, *Games and Economic Behavior* 16, 280–302.
- Shleifer, Andrei, and Robert Vishny, 1989, Management entrenchment: The case of manager-specific investments, *Journal of Financial Economics* 25, 123–139.
- Spencer Stuart, 1997, Board Index: Board trends and practices at S&P 500 corporations (Spencer-Stuart, Chicago).
- Subrahmanyam, Vijaya, Nanda Rangan, and Stuart Rosenstein, 1997, The role of outside directors in bank mergers, *Financial Management* 26, 23–36.
- Warther, Vincent, 1998, Board effectiveness and board dissent: A model of the board's relationship to management and shareholders, *Journal of Corporate Finance* 4, 53–70.
- Weisbach, Michael, 1988, Outside directors and CEO turnover, *Journal of Financial Economics* 20, 431–460.
- Yermack, David, 1996, Higher market valuation of companies with a small board of directors, *Journal of Financial Economics* 40, 185–211.