

团队生产防联盟激励机制设计

代建生*

摘要: 本文运用“核”的基本思想,通过引入防联盟激励约束,构建了团队生产防联盟相互激励模型,并刻画了模型解的一阶最优条件及其解域。在团队激励机制设计中,只有当所有团队成员参与极大联盟所得收益大于参与一切可能的子联盟所得收益时,激励机制才可能是防联盟的。运用防联盟激励模型对一些经济现象给出了不同于传统理论的解释,比如:薪酬结构差异、效率工资以及内部人控制现象等,并将上述经济现象背后的原因部分归结为团队成员的谈判能力和结盟的可能性。研究指出,在团队激励机制设计中应重点关注某些可能结盟的工作团队。本研究可用于指导三人以上团队在激励成员积极投入的同时防止部分成员结盟。

关键词: 团队生产 防联盟 激励机制 核

一、引言

在团队生产中,团队成员的投入对联合产出的贡献难以准确计量,容易诱发团队成员搭便车等道德风险问题,致使团队合作的效率低下(Alchian and Demsetz, 1972)。通过引入委托人打破预算平衡,或让某一团队成员成为反分享者,有助于维系团队的效率(Eswaran and Kotwal, 1984; Holmstrom, 1982);借助集体声誉补贴团队生产中的搭便车带来的效率损失,能有效缓解预算平衡约束和激励相容约束之间的矛盾(李金波等, 2010)。

代建生和孟卫东(2010)通过引入团队福利函数考察了合伙型自主工作团队的激励问题,构建了相互激励模型,他们指出机制设计者要吸引成员参与团队生产,给成员的补偿支付不能小于其保留效用。与只有两个参与者的情形有所不同,在三人以上的团队中,团队成员即使不参与极大联盟(grand coalition),仍可结成子联盟(sub-coalition)进行合作,各成员所得收益将大于其保留效用,因为把联盟视为一体与对方谈判可提高联盟成员的支付(Hart and Kurz, 1983)。因此,在团队激励机制设计中,要吸引团队成员参与极大联盟,仅仅满足个体参与约束是不够的;只有当参与极大联盟所得收益大于参与一切可能的子联盟所得收益时,参与极大联盟对所有成员而言才是有利可图的。

在经济学文献中,术语 coalition-proof 和 collusion-proof 既有区别又有联系,前者通常译为“防(抗)联盟”,它从经典的合作博弈理论演化而来,运用特征函数对联盟(coalition)结构进行刻画,不考虑联盟成员的策略选择对其他成员收益的外部性(Maskin, 2011);后者通常译为“防(抗)合谋”,强调非合作博弈的基础,充分考虑合谋参与者的策略选择对其他参与者收益

* 代建生,昆明理工大学管理与经济学院,邮政编码:650093,电子信箱:jiansheng.dai@163.com。

本文的研究受到中央高校基本科研业务费项目“防共谋团队生产理论及其在研发联盟激励机制设计中的应用研究”(项目编号:CDJSK10 02 01)和中国博士后科学基金资助项目“有限理性下产学研合作研发防合谋激励机制设计”(项目编号:2012M511909)的资助。感谢匿名审稿人提出的宝贵建议,当然文责自负。

的外部性影响。前者主要出现在抽象的合作博弈理论(Moreno and Wooders,1996)、社会选择理论(Jackson and Sonnenschein,2007;Sengupta and Sengupta,1996)、机制设计及其实施理论(Boylan,1998),以及市场理论(Delgado and Moreno,2004;Kamishiro,2011)等文献中,后者更多地出现在劳动合约理论(Bertomeu,2007;Brown and Wolfstetter,1989)、规制理论(Laffont and Martimort,1997)、审计理论(Baiman, et al.,1991)和拍卖理论(Che and Kim,2009)等文献中。尽管存在这些差别,两者却有着内在的一致性,一是合谋(collusion)的参与者事实上组成极大联盟中的子联盟。二是防止联盟(合谋)的逻辑是一致的:要防止极大联盟中的部分成员结成子联盟,成员所得支付应在极大联盟博弈的核中,使得部分成员结盟无利可图(Boylan,1998;Delgado and Moreno,2004;Jackson and Sonnenschein,2007;Kamishiro,2011;Moreno and Wooders,1996;Sengupta and Sengupta,1996);而要防止经济参与者的合谋,合约设计者可设计一个契约使得代理人从中得到的收益不少于合谋收益,致使代理人没有进行合谋的积极性(Brown and Wolfstetter,1989;Che and Kim,2009;Laffont and Martimort,1997;陈志俊、邱敬渊,2003)。

在信息不对称的多边合约中,如果合约不是联盟激励相容的,合约将不可实施(Forges, et al.,2002;Green and Laffont,1979)。由于代建生和孟卫东(2010)没有考虑联盟问题,他们设计的激励机制不是联盟激励相容的,其结果只适用于两个参与者的情形或不存在结盟可能性的 n 个参与者的情形,本文致力于将他们的结果推广到防联盟的情形。

在现有的防合谋文献中,由于强调合谋的外部性致使需要处理的问题极其复杂,这类文献通常仅对三个参与者(委托人、监督者和代理人)的情形进行讨论(Brown and Wolfstetter,1989;Che and Kim,2009;Grimaud, et al.,2003),而难以对更一般的情形进行处理。不仅如此,在这类文献中,要考虑的合谋通常发生在特定的两个成员之间,比如代理人和监督者(Grimaud, et al.,2003;董志强、严太华,2007;蒋神州,2011),或者两个代理人之间(Che and Kim,2009;陈志俊、邱敬渊,2003),这又使问题简化了。

在合伙型自主工作团队中,所有成员的地位是平等的,合约通过谈判达成而不是由某个委托人单方面提出,并且合约是在事前达成,不清楚事后哪几个成员将组建联盟,所有成员之间的联盟(合谋)都是可能的。对于一般的 n 个参与者的情形,如果将单个经济人视为一个子联盟,共需处理 $2^n - 1$ 个联盟,这将使问题极其复杂。尽管如此,如果不考虑联盟的外部性问题,则可借助合作博弈理论中“核”的思想来讨论防联盟问题。特别地,将各种子联盟的收益用特征函数来表征,将极大地简化对问题的处理。

基于此,本文暂不考虑联盟的外部性问题,而是运用“核”的概念来讨论防联盟问题,并借助特征函数将联盟激励约束转化为防联盟约束(集体理性约束),力求在一个统一的模型中综合考察团队生产中的防联盟问题和激励问题,其基本思路是:团队的机制设计者设计一个激励合约,使得团队成员的策略选择符合机制设计者的利益,且确保团队成员的支付位于团队博弈的核中,使得团队成员没有动机结成子联盟。技术上的简化处理基于以下理由:一是假定子联盟一旦结成,联盟契约就具有约束力,因而机制设计者可将子联盟作为单个经济参与者来处理(Green and Laffont,1979);二是假定子联盟成员的策略选择对联盟之外的成员收益没有外部性(Maskin,2011)。为了讨论的方便,本文假定极大联盟的核总是存在的,这是合作博弈理论文献的惯常做法(Maskin,2011)。

本文的结构安排如下:第二部分描述基本模型;第三部分引入防联盟机制;第四部分考察防联盟激励模型解的性质及解域;第五部分运用防联盟激励模型解释部分经济现象;最后是结论及研究展望。

二、基本模型及假设

在联合生产中,不能仅仅通过对产出 π 的观察完美地推导 e ,这是因为产出不仅受到努力水平的影响,还与随机因素有关。给定 e ,团队总产出 π 服从随机分布 $F(\pi|e)$ 。

不失一般性,约定条件密度函数 $f(\pi|e) > 0, \forall \pi \in [\underline{\pi}, \bar{\pi}], \forall e \in E$ 。这个假定表示对事先给定的努力水平组合 e ,由于随机因素的影响,位于定义域中的任意产出都有可能实现。

团队本身的行动是由一个虚拟的计划者或说由其所有成员通过谈判共同决定,通过制定分配规则来激励所欲实施的努力(代建生、孟卫东,2010)。团队追求团队福利函数 $W(u)$ 的极大化,本文考虑一个特殊的福利函数: $W(u) = \sum_{i \in N} \tau_i h_i(u_i)$, 其中 $h_i(\cdot)$ 为转换函数,它是一个关于 u_i 的严格递增的函数,即 $h'_i(\cdot) > 0$ 且 $h''_i(\cdot) \leq 0$, 而 $\tau_i \geq 0$ 是成员 i 的谈判力因子。不失一般性,将谈判力因子作归一化处理。团队成员选择是否参与团队,如果参与,从策略空间中选择最优策略。

由于信息不对称,团队需要通过内部分配规则(w, π)在确保成员参与团队生产的同时激励成员加大投入,以极大化团队福利。团队面临以下问题:

团队追求团队福利的极大化,团队福利用团队福利函数 $W(u)$ 来表征,一些特殊的福利函数如独裁解、均等解、功利解、纳什谈判解及 K-S 解等(代建生、孟卫东,2010)。在模型(1)中,约束 IC 是激励相容约束,由于行动的不可观察,团队成员基于自身利益从其策略空间中选择最优策略。约束 BB 为预算约束,团队成员获取的收益之和不大于团队总产出。我们关注的重点是 IR 约束,它是吸引成员参与团队合作的必要条件,在三人及其以上的团队中这一约束是不充分的,本文的主要工作就是拓展这一约束条件。

三、防联盟分配机制

(一) 基本概念

92

即 $u = (u_1, \dots, u_n) \in R^n$, 其子联盟 S 的支付向量用 u^S 来表示, $u^S = (u_i)_{i \in S} \in R^S$ 。

定义 1 (Mas - Collet, et al., 2005): 对于 $S \subseteq N$, 一个非空闭集 $U^S \subset R^S$ 是效用可行集, 如果它满足: 对任意 $u^S \in U^S$, 有: $u^S - R_+^S \in U^S$ 。

定义 2 (Mas - Collet, et al., 2005): 团队 N 的合作博弈用特征函数 (N, V) 来表示, 这个博弈指定以下分配规则: 对团队 N 中任意联盟 S , 分配一个效用可行集 $V(S) \subset R^S$ 。

定义 3 (Mas - Collet, et al., 2005): 以特征形式表示的效用可转移博弈, 记为 (N, v) , 其中 N 表示参与人的集合, 而 v 是一个函数, 它赋予 N 中任意联盟 S 一个数值 $v(S)$ 。

(二) 防联盟分配机制

1. 产出确定情形下的防联盟分配

在合作博弈 (N, V) 中, 联盟 S 为核非空联盟, 如果存在一个可行分配 $u^S \in R^S$, 它不能被 S 的任意子联盟所优越; 否则, 联盟 S 为空核联盟。所有核非空联盟的集合记为 $P_R(N)$, 特别地, 约定 $\{i\} \in P_R(N)$ 且 $N \in P_R(N)$, 即单个局中人的联盟和极大联盟均属于核非空联盟。在本文中, 如没有特别说明, 当提到联盟 S 时, 总是假定 $S \in P_R(N)$ 。

定义 4: 对于 $S \in P_R(N)$ 以及真子集 $T \in P_R(N)$, 如果不存在那样一个集合 R , 有: $T \subset R \subset S$ 且 $R \in P_R(N)$, 则说 T 是 S 的极大真子联盟, S 的极大真子联盟的全体记为 $Q_R(S)$ 。

为了讨论的方便, 对于博弈 (N, V) , 约定 $V(\{i\}) = 0$, 对任意 i 成立, 即成员 i 不与其他成员组建联盟时的保留效用为零。对效用可行集 $U^S \subset R^S$, 如果 $u^S \in R_+^S - U^S$, 那么 u^S 在总体上优于集 U^S , 即 $u^S > U^S$ 。一切 U^S 的总体优于集记为 $U^{>S}$ 。对于 $U^{>S}$, 不能保证对一切 $u'^S \in R_+^S - U^S$ 及 $u^S \in U^S$, 有 $u'^S \geq u^S$, 在本质上, 总体优于集是不可实现的分配的集合。类似地, $V(S)$ 的总体优于集记为 $V^>(S)$, 即 $V^>(S) = R_+^S - V(S)$ 。

对于 $u^S \in U^S$, 且 $u^S \geq 0$, 称分配 u^S 位于 U^S 的帕累托效率前沿, 如果不存在另一可行分配 $u'^S \in U^S$, 有 $u'^S \geq u^S$, 且其中至少一个严格不等式成立; U^S 的帕累托效率前沿用记号 $U^{\wedge S}$ 表示。

令 $U^{\geq S} = U^{>S} \cup U^{\wedge S}$, 后文提及 $U^{\geq S}$ 为总体至少一样优集。

命题 1: 给定博弈 (N, V) 及其可行分配 $u \in R^n$, 如果 $u^S \in V^{\geq}(S)$, $\forall S \in P_R(N)/N$, 那么 u 具有防联盟性, 或说 u 是防联盟的。

根据命题 1, 要判定分配 u 是否防联盟, 只需考察 $P_R(N)$ 中的子联盟, 而不必关注其他子联盟。

推论 1: 给定博弈 (N, v) 及可行分配 $u \in R^n$, 如果 $\sum_{i \in S} u_i^S \geq v(S)$, $\forall S \in P_R(N)/N$, 那么 u 是防联盟的。

2. 产出随机情形下的防联盟分配

当产出 π 随机时, 设其分布函数为 $F(\pi)$, 记 $Eu_i(\cdot) = \int u_i(\cdot) dF(\pi)$ 。

当要强调分布函数 $F(\pi)$ 与策略 e 的联系时, 记 $E_e u_i(\cdot) = \int u_i(w_i(\pi), g_i(e_i)) dF(\pi|e)$, 其中 $F(\pi|e)$ 是与策略 e 相对应的产出条件分布函数。

给定产出 π , 用 $V_{|\pi}(N) \subset R^n$ 来表示极大联盟的效用可行集。设在产出 π 给定下的一个分配 $w(\pi)$, 记 $u(w(\pi)) = (u_1(w_1(\pi)), \dots, u_n(w_n(\pi)))$, 其效用可行集为: $V_{|\pi}(N) = \{u(w(\pi)) : \sum_{i \in N} w_i(\pi) \leq \pi\}$ 。

定义 5: 以特征函数 $(N, V, F(\pi))$ 表示的随机博弈是指团队 N 的产出分配博弈: 任给一个产出 π , 对极大联盟 N , 赋予一个效用可行集 $V_{|\pi}(N)$; 而对 N 中任意真子联盟 S , 赋予一个效用

可行集 $V(S)$ 。

命题2:对于以特征形式给出的随机博弈 $(N, V, F(\pi))$, 其产出 π 服从随机分布 $F(\pi)$, 如果对 N 的一切真子联盟 $S \in P_R(N)/N$, 有: $Eu^S \in V^\geq(S)$, 那么相应的可行分配 $w(\pi) \in R^n$ 是防联盟的。

命题2说明的是, S 中的成员选择参与极大联盟 N 所得支付在期望意义上至少与不参与联盟 N 而组成联盟 S 所得支付一样好时, N 中的成员没有必要形成联盟 S 。

推论2:对于效用可转移博弈 $(N, v, F(\pi))$, 称可行分配 $w(\pi) \in R^n$ 是防联盟的, 如果对 N 的一切真子集 $S \in P_R(N)/N$, 有 $\sum_{i \in S} Eu_i(w_i) \geq v(S)$ 且 $Eu_i^S(w_i) \geq 0$ 。

四、模型扩展及分析

(一) 团队相互激励模型的扩展

本部分将运用防联盟支付取代模型(1)中的保留效用, 并将模型(1)推广到防联盟的情形。

令 $E_e V(N) = \{u(w(\pi), g(e)) : u(w(\pi), g(e)) \in V_{|\pi}(N), \forall \pi\}$ 。运用这一记号, 将防联盟激励模型表示如下:

$$\begin{cases} \max_{(w,e)} \int W(u) f(\pi|e) d\pi \\ \text{s. t. (CP)} Eu^S \in V^\geq(S), \forall S \in P_R(N)/N \\ \quad \text{(IC)} e_i \in \arg \max_{e'_i \in E_i} \int u_i(w_i(\pi), e'_i) f(\pi|e'_i, e_{-i}) d\pi, \forall i \in N \\ \quad \text{(BB)} u \in E_e V(N) \end{cases} \quad (2)$$

其中, CP 约束是防联盟约束, 这一约束要求团队产出的分配满足防联盟性。模型(1)中的 IR 约束是团队为了阻止合作博弈中的成员个体通过单独行动来摆脱团队而施加的约束; 与之有所不同, CP 约束是团队为了阻止成员个体单独或联合行动来摆脱团队而施加的约束, 在这个意义上, CP 约束是 IR 约束的一个自然扩展。在防联盟激励模型中, 当 $Q_R(N) = \{\{i\}, i \in N\}$ 时, 防联盟激励模型退化为非联盟情形下的激励模型, 即模型(1)。如果机制设计者知道联盟结构, 那么机制设计者可以将联盟视为单个代理人对待。但由于合约只能在事前提供, 此时机制设计者对哪些成员可能结成联盟是不清楚的, 因而机制设计者必须考虑到所有可能的子联盟。

令 $E_e \text{core}(N) = \{u : Eu^S \in V^\geq(S), \forall S \in P_R(N)/N, \text{ and } u(w(\pi), g(e)) \in V_{|\pi}(N), \forall \pi\}$ 。利用这一记号, 重写模型(2)如下:

$$\begin{cases} \max_{(w,e)} \int W(u) dF(\pi|e) \\ \text{s. t. (CP - BB)} u \in E_e \text{core}(N) \\ \quad \text{(IC)} e_i \in \arg \max_{e'_i \in E_i} \int u_i(w_i(\pi), e'_i) dF(\pi|e'_i, e_{-i}), \forall i \in N \end{cases} \quad (3)$$

模型(3)用约束 CP - BB 替换了模型(2)中的约束 CP 和 BB。CP - BB 约束是对效用可行域施加的限制, 要求团队的分配不仅具有可行性, 而且要具有防联盟性。在本质上, 这一约束条件刻画的是模型(2)解的可行域, 即团队博弈 $(N, V, F(\pi))$ 的核。

(二) 防联盟激励模型的一阶最优条件

对任意现实的真子联盟 S 及其分配 $u \in V^\wedge(S)$, 当 u 连续时, $V^\wedge(S)$ 上的点可用 s 维空间中一连续曲面来刻画。特别地, 如果 u 关于 w 拟凹, 那么 $V^\wedge(S)$ 是一凹向原点的 $s - 1$ 维曲面。

为了讨论方便,用方程 $H_S(\cdot) = 0$ 来描述联盟 S 的帕累托效率前沿, $H_S(\cdot) \leq 0$ 表示联盟 S 所能实现的效用可行域,而 $H_S(\cdot) > 0$ 表示联盟 S 所不能实现的效用可行域。使用这些符号,重写模型(3) 如下:

$$\begin{cases} \max_{(w,e)} \int W(u)f(\pi|e)d\pi \\ \text{s. t. (BB)} H_N(u) \leq 0 \\ \quad \text{(CP)} \int u_i(w_i(\pi), e_i)f(\pi|e)d\pi \geq u_i^S, \text{ where } H_S(u^S) \geq 0, \forall S \in P_R(N)/N \\ \quad \text{(IC)} e_i \in \arg \max_{e'_i \in E_i} \int u_i(w_i(\pi), e'_i)dF(\pi|e'_i, e_{-i}), \forall i \in N \end{cases} \quad (3')$$

令 $\lambda_S, \xi(\pi)$ 和 $\mu_{e'_i}$ 分别为相应的约束 CP、BB 和 IC 的拉格朗日乘子,注意到 $H_N(u) \leq 0$ 等价于 $\sum_{i \in N} w_i(\pi) \leq \pi$, 则模型(3') 关于 $w_i(\pi)$ 的一阶最优条件为:

$$\frac{\xi(\pi)}{v'_i(\cdot)f(\pi|e)} = \tau_i h'_i(\cdot) + \sum_{i \in S, S \in P_R(N)/N} \lambda_S + \sum_{e'_i \neq e_i} \mu_{e'_i} \left(1 - \frac{f(\pi|e'_i, e_{-i})}{f(\pi|e_i, e_{-i})} \right), \forall i \in N \quad (4)$$

当策略集 E 连续时,其 IC 约束的一阶条件为:

$$\int [v_i(w_i(\pi))] f_{e_i}(\pi|e) d\pi - g'_i(e_i) = 0。$$

使用上式替换 IC 约束,令 μ_i 为其乘子,则模型(3') 在连续情形下的一阶最优条件为:

$$\frac{\xi(\pi)}{v'_i(\cdot)f(\pi|e)} = \tau_i h'_i(\cdot) + \sum_{i \in S, S \in P_R(N)/N} \lambda_S + \mu_i \frac{f_{e_i}(\pi|e)}{f(\pi|e)}, \forall i \in N \quad (4')$$

命题3:在模型(3) 中,如果 $\tau_i = 0$,那么在式(4) 或式(4') 中,必然存在某个 S , 有 $\lambda_S > 0$, 其中 $i \in S \subset N$ 。

引理1:给定 e , 在式(4) 或式(4') 中,至少存在一个测度不为零的空间 Ω , 有 $\xi(\pi) > 0$, $\forall \pi \in \Omega$ 。

引理1 和命题3 的证明过程见附录1 和附录2。

根据命题3, 如果某个成员的谈判力因子等于零,即他相对于其他团队成员没有任何谈判能力,那么在该成员可能参加的所有子联盟中,至少有一个子联盟的防联盟约束是起作用的,即该子联盟没有从极大联盟的合作中分配到团队生产的净剩余。在经典委托代理模型中,其最优解要求 IR 约束起作用,注意到代理人毫无谈判能力,这表明命题3 是经典代理模型相关结论的自然推广。

(三) 防联盟激励模型的解域

定义6:对于集 $core(N)$, 如果 u 满足 $u \in core(N)$, 且 $u + R_+^n / \{0\} \notin core(N)$, 则称 u 位于 $core(N)$ 的上边界, 一切那样的 u 的集合构成 $core(N)$ 的上边界; 如果 $u \in core(N)$, 且 $u - R_+^n / \{0\} \notin core(N)$, 则称 u 位于 $core(N)$ 的下边界, 一切那样的 u 的集合构成 $core(N)$ 的下边界。记:

$$core^\wedge(N) = \{u: u \in V^\wedge(N), u^S \in V^\geq(S), \text{ for any } S \subset N\},$$

$$core^\vee(N) = \{u: u \in V(N); u \in U^\vee(T), \exists T \subset Q_R(N)\}。$$

其中, $U^\vee(T) = \{u: u \in V(N), u^T \in V^\wedge(T), u^{N-T} \in V^\wedge(N-T), T \in Q_R(N), u^S \in V^\geq(S), \forall S \in P_R(N)/N\}。$

命题4: 集合 $core(N)$ 具有以下性质: (i) $core(N)$ 有界; (ii) $core(N)$ 的上界为

$core^{\wedge}(N)$, 下界为 $core^{\vee}(N)$ 。

证明过程见附录3。

设 $\sum_k^p \tau_{i_k} = 1$ 且 $\tau_{i_k} > 0$, 其中 $p \geq 1$ 。为了记号的简化, 令 $\{i\} = \{i_1, i_2, \dots, i_p\}$, $\{j\} = \{j_1, j_2, \dots, j_{n-p}\}$ 。

很显然, 集合 $\{j\}$ 中的成员毫无谈判能力, 由其成员组成的集合称为代理人集团, 而 $\{i\}$ 中成员的谈判力因子均大于零, 该成员集合称为委托人集团。本文将由 p 个委托人集团的成员参与的团队生产模型称为 p -委托人- $(n-p)$ -代理人模型。特别地, 当 $p = 1$ 时, 团队中某个成员拥有完全的谈判能力, 这类模型被称为完全委托代理模型。

用 N 表示极大联盟, P 表示委托人集团, 而 $N-P$ 为纯粹代理人集团。对纯粹代理人集团而言, 存在以下两种情况: $N-P \in P_R(N)$ 以及 $N-P \notin P_R(N)$ 。特别地, 当 $N-P \notin P_R(N)$ 时, 我们用 $Q_R(N-P)$ 表示代理人集团的极大真子联盟。令:

$$V_{N-P}^{\wedge}(T) = \left\{ u^{N-P}; u^{N-P} \in V(N-P), u^T \in V^{\wedge}(T), u^{N-P-T} \in V^{\wedge}(N-P-T), \right. \\ \left. T \in Q_R(N-P); u^S \in V^{\geq}(S), \forall S \in N-P \right\},$$

$$V^{\wedge}(N-P) = \bigcup_{T_i \in Q_R(N-P)} V_{N-P}^{\wedge}(T_i)。$$

一般地, 委托人集团和代理人集团内部以及两类集团之间的部分成员都有可能结成子联盟, 如果两类集团的成员没有结盟的可能性, 则有以下结论:

命题5: 对于 p -委托人- $(n-p)$ -代理人防联盟激励模型(3), 如果不存在委托人与代理人之间联盟的可能性, 那么其任意最优解 u^* 必然满足: $Eu^{(N-P)*} \in V^{\wedge}(N-P)$ 。

证明过程见附录4。

纯粹代理人集团不参与团队净剩余的分配, 他们从团队中获得的最大收益, 与完全由这一集团组建的子联盟所能实现的最大收益(即该子联盟的保留收益)等价。由于代理人集团没有谈判能力, 他们不能参与合约的制定过程, 只是被动地接受或者拒绝由委托人集团单方面提供的合约, 而且, 代理人集团和委托人集团之间也不存在部分成员跨集团联盟的可能, 在这种情形下, 委托人集团必然提供那样的合约安排, 使得代理人集团刚好获得保留收益, 这对委托人集团是最优的。

团队的分配规则由团队成员通过谈判共同决定, 这是一般的情形。事实上, 真正参与分配规则制定的是拥有谈判力的成员, 对 p -委托人- $(n-p)$ -代理人模型而言, 就是委托人集团参与了规则的决定, 而代理人集团, 只是被动地接受分配规则, 就像经典代理模型中代理人是规则的被动接受者一样。

推论3: 设 u^* 是完全委托代理模型的最优解, 则有以下结论: $Eu^{(N-1)*} \in V^{\wedge}(N-1)$, 其中“1”代表委托人, “ $N-1$ ”表示代理人集团。

推论3表明代理人集团不能参与合作净剩余的分配。在经典代理模型中, 代理人得到的效用碰巧等于其保留效用, 推论3将这一结论推广到存在多个代理人的情形。由命题5, 如果代理人集团和委托人集团之间不存在部分成员跨集团联盟的可能性, 那么代理人集团只能得到保留收益。在完全委托代理模型中, 只存在一个拥有完全谈判能力的委托人, 他没有与代理人集团中的任何代理人结盟的激励, 因为让代理人集团刚好得到保留收益是委托人的最优策略。

五、模型运用

(一) 企业不同员工薪酬结构的差异

在企业中普通员工通常只能获得薪酬、福利或奖金(经济激励)而不能获得股权激励(企业剩余分享), 而股东及高层管理人员则可获得包括股权激励在内的一些利益分配, 参与企业

净剩余的分享。拥有特殊技能的高技术人才也能参与企业净剩余的分配,尤其是在高新技术企业中(方阳春、姚先国,2007)。

用本文的术语,在企业这个团队中,股东及高级管理人员可视为委托人,拥有特殊技能的高级技术人才也可视为委托人,他们共同瓜分企业的净剩余。而普通员工,由于没有任何谈判力,则不能参与企业净剩余的分配。

卢周来(2009)认为,企业要素拥有者在企业内的谈判能力取决于要素的重置成本,而这又最终取决于要素市场的稀缺程度。普通员工只具有一般的劳动资源,这类资源从市场供需来看并不稀缺,其资源几乎处在完全竞争的市场中,因而面对大股东和高级管理层几乎没有任何谈判能力。在上市公司中,甚至小股东也没有什么谈判力,他们只能追求资金的市场平均收益,因为:(1)大股东与公司高层存在合谋的可能性;(2)更重要的是小股东的资源不那么稀缺。拥有特殊技能的人才则不同,他们拥有相对稀缺的智力资源,尤其是CEO等高管人才,对经营的企业拥有其他人所不具有的独特的能力,难以替代,因而具有较强的谈判能力。企业需要运用包括股权激励在内的一些手段吸引他们留在企业这个团队中,并允许他们分享企业创造的合作剩余。

(二)效率工资现象^①

这里所说的效率工资现象,是指企业一般员工可获得高于由市场出清决定的工资水平的经济现象。根据卢周来(2009),既然普通员工拥有的劳动资源处在完全竞争市场中,其重置成本就是员工不参与企业可获得的保留效用,但这与企业往往会给予员工高于由市场确定的工资水平的薪酬相矛盾,这是他所不能解释的现象。在他建立的模型中,仅仅考察了两个成员的纳什讨价还价问题,没有为讨论成员之间的子联盟问题留下任何空间。根据本文的研究,尽管普通员工面对企业高管和股东没有什么谈判能力,但他们却可能结成联盟进行讨价还价,比如工会。

企业可通过解聘的方式促使单个或少数员工仅获得由市场出清决定的工资水平,但不能对所有或数量众多的员工采取这种方式,尤其是某些员工团队^②,这些团队在离开该企业后,团队整体或部分成员仍可加盟到其他企业,并能创造比团队成员单干所能创造的效益之和更高的效益。因此,要防止此类团队集体跳槽,就必须给予他们这样的薪金水平,使得从团队总体来看,跳槽不是最优选择,即薪金水平满足团队的集体理性。

普通员工获得高于市场确定的工资水平并不表明他们拥有了更高的谈判能力,这只是因为一些员工可能结盟,而结盟后形成的子联盟整体的防联盟支付大于子联盟中各个成员保留效用之和。企业为了防止他们结盟,给予他们高于市场工资水平的薪金收入。尽管如此,所有没有谈判能力的代理人的结盟,他们从整体上所能获得的最高收益不会大于当他们离开本企业后所可能获得的最大收益。否则,委托人集团可以降低部分成员的薪酬水平,在能防止他们联盟的同时,使委托人集团获得更大的收益。

特别地,并非任意几个普通员工都具有结成现实的子联盟的可能性。与企业中一些工作团队,如研发团队、销售团队或财务团队等有所不同,彼此关系不大的员工未必能结成有实际意义的联盟,即这样的团队可能不属于真子联盟集合 $P_R(N)$,因为这样的数个员工跳槽后联合

^①效率工资理论从其他角度解释了这一经济现象,其中一种理论认为:企业通过支付员工高于由市场出清决定的工资水平,将增大员工离职的机会成本,减少离职率。这种理论与我们的观点相近,但我们是从防联盟这一角度对此进行讨论的。

^②企业中实际存在的一些团队,如:设计团队、销售团队、律师团队、财务团队等等,很可能属于极大联盟的真子联盟集合 $P_R(N)$,在防联盟机制设计中应成为重点关注的对象。

创造的收益未必大于他们的保留效用之和。而且,这些员工由于种种现实的原因,比如相互不认识,甚至完全没有结盟的可能性。

(三) 内部人控制

在现代企业中,所有权与经营权高度分离,经营者控制了公司的筹资权、投资权、人事权等经营权,因所有者与经营者在利益方面的冲突,经营者可能采取不利于所有者利益的行动,而所有者对经营者却难以施加有效的控制,这种现象被称为内部人控制。在国有企业中,内部人控制现象尤为明显,国有企业内部人可以从企业投资行为中获得隐性的私人收益,或实施利润转移行为以增进自身利益(钱雪松、孔东民,2012)。

传统理论认为,剩余索取权和控制权的不相匹配是导致内部人控制的主要原因之一。比如,在国有企业中,由于拥有剩余索取权的所有者的主体缺失,使得所有者失去了对企业的实际控制权;而企业的经营者和管理人员,凭借对企业的经营决策获得了对企业的实际控制权。由于拥有剩余索取权的所有者没能实际控制企业,且对企业经营者的监督乏力,滋长了经营者对企业的内部人控制问题。

除了这些理由,我们认为以下原因也滋长了内部人控制现象:一是因为股权分散,每个股东相对于企业管理层的谈判力很小,并且股东之间不易形成联盟来对抗管理层,这一点在国有企业中表现尤为明显;另一方面,相对于股东,企业管理层因为其特殊的资源(比如管理能力)而拥有较强的谈判力,不仅如此,企业管理层还可能达成默契共谋,并形成利益集团,以至于将企业置于管理层控制之下。

联盟不仅可能发生在普通员工之间,企业高层管理人员之间,或者股东之间,也可能发生在员工与高层管理人员之间,或者股东与高层管理人员之间。无论是哪种结盟,其目的都在于获取更高的收益。

六、结束语

运用“核”的基本思想,并借助特征函数这一工具,本文将相互激励模型推广到防联盟情形。运用防联盟激励机制,在激励团队成员增大投入,减少因道德风险导致的“搭便车”问题的同时,还能有效地防止团队中部分成员结盟的可能性,使得激励机制能有效地实施。

参与团队分配规则制定的是拥有谈判力的成员,没有谈判力的成员是规则的被动接受者,在他们所参加的所有可能的子联盟中,至少有一个子联盟不能参与团队净剩余的分配。当团队中部分成员能结成现实的联盟时,团队激励机制设计必须是防联盟的,其管理意义是,在防联盟机制设计中应重点关注某些工作团队,如设计团队、销售团队、律师团队等。

运用防联盟激励模型对一些经济现象做出了不同于传统理论的解释,如薪酬结构差异、效率工资及内部人控制现象等,我们将上述经济现象背后的原因部分归结为团队成员的谈判能力和结盟的可能性。

最后,我们认为,有以下几个方向值得进一步研究:第一,在实证方面,本文的研究缺乏现实的直接证据,下一步可考虑加强实证研究,以检验文中结论;第二,在应用方面,将防合谋激励模型用于解释更多的社会经济现象,或者对文中提到的经济现象作更为深入透彻的研究;第三,在理论方面,本文暂未考虑联盟的外部性问题,引入外部性尽管面临技术复杂性难题,但我们相信这是一个有趣的研究方向。

附录 1:引理 1 证明

证明:只证明离散情形,连续情形可类似证明。首先注意到,由约束极值理论,有: $\xi(\pi) \geq 0, \forall \pi \in [\underline{\pi}, \bar{\pi}]$,及 $\lambda_s \geq 0, \forall S \in P_k(N)/N$ 。假设 $\xi(\pi) = 0$, almost everywhere, 必有:

$$\tau_i h'_i(\cdot) + \sum_{i \in S, S \in P_R(N)/N} \lambda_S + \sum_{e'_i \neq e_i} \mu_{e'_i} \left(1 - \frac{f(\pi | e'_i, e_{-i})}{f(\pi | e_i, e_{-i})} \right) = 0, \forall i \in N.$$

这要求 $\sum_{e'_i \neq e_i} \mu_{e'_i} \left(1 - \frac{f(\pi | e'_i, e_{-i})}{f(\pi | e_i, e_{-i})} \right) < 0$, almost everywhere, 对某个满足 $\tau_i > 0$ 的成员 i 。但这是不可能的,

因为 $\int f(\pi | e_i, e_{-i}) d\pi = \int f(\pi | e'_i, e_{-i}) d\pi, \forall e$, 从而导致矛盾。证毕。

附录 2: 命题 3 证明

证明: 同样只证明离散情形, 连续情形可类似证明。

假设 $\lambda_S = 0, \forall S$, 且 $i \in S$ 。由式(4) 或式(4'), 有:

$$\sum_{e'_i \neq e_i} \mu_{e'_i} \left(1 - \frac{f(\pi | e'_i, e_{-i})}{f(\pi | e_i, e_{-i})} \right) \geq 0, \forall \pi \in [\underline{\pi}, \bar{\pi}].$$

但根据引理 1, 对一个测度不为零的空间 Ω , 又有:

$$\sum_{e'_i \neq e_i} \mu_{e'_i} \left(1 - \frac{f(\pi | e'_i, e_{-i})}{f(\pi | e_i, e_{-i})} \right) > 0, \forall \pi \in \Omega.$$

这与 $\int f(\pi | e_i, e_{-i}) d\pi = \int f(\pi | e'_i, e_{-i}) d\pi$ 对一切 e 成立相矛盾。证毕。

附录 3: 命题 4 证明

证明: (i) 由定义可完成证明。

(ii): 容易表明 $core(N)$ 的上界为 $core^\wedge(N)$ 。关于其下界, 首先注意到对任意 $u \in core^\vee(N)$, 有: $u - R_+^n / \{0\} \notin core(N)$, 就是说 $core^\vee(N)$ 位于 $core(N)$ 的下边界上。下证对任意位于 $core(N)$ 的下边界的 u , 满足 $u \in core^\vee(N)$ 。如若不然, 则存在 u , 满足 $u \in V(N)$, 且对某一 $T \in Q_R(N)$, 要么 $u^T \in V^>(T)$, 要么 $u^{N-T} \in V^>(N-T)$, 如此, 则可找到某一 u' , 满足 $u' \leq u$, 且 $u' \in core(N)$ 。证毕。

附录 4: 命题 5 证明

证明: 假设 $Eu^{(N-P)*} \notin V^\wedge(N-P)$, 那么 $Eu^{(N-P)*} \in V^>(N-P)$ 。在 $N-P$ 中一定存在某个 i , 有 $Eu_i(w_i^*, e_i^*) > u_i^0$ 。考虑一个满足 IR 约束的分配 w :

$$w_i < w_i^*, i \in N-P, \text{使得 } u_i(w_i, e_i^*) + \delta = u_i(w_i^*, e_i^*), \forall \pi \in [\underline{\pi}, \bar{\pi}];$$

$$w_j > w_j^*, j \in P, \text{使得 } u_j(w_j, e_j^*) = u_j(w_j^*, e_j^*) + \varepsilon, \forall \pi \in [\underline{\pi}, \bar{\pi}];$$

$$w_j - w_j^* \leq w_i^* - w_i, \forall \pi \in [\underline{\pi}, \bar{\pi}];$$

$$w_k = w_k^*, \forall k \in N \text{ 且 } k \neq i, j.$$

既然 u 在定义域内连续有界且关于 w 严格递增, 给定任意 δ , 一定存在相应的 ε 使上面的(不)等式成立。设 u 是与 w 相对应的效用向量, u 满足 IC 约束和 w 满足 BB 约束是明显的。且只要 δ 足够小, 就可保证 u 满足 CP 约束。注意到 $W(\cdot)$ 是关于 u^p 的严格增函数, 而 u 是关于 w 的严格增函数, 则必有 $W(u(w)) > W(u(w^*))$, 这与假设相矛盾。证毕。

参考文献:

1. 陈志俊、邱敬渊, 2003:《分而治之: 防范合谋的不对称机制》,《经济学(季刊)》第 1 期。
2. 代建生、孟卫东, 2010:《团队生产中的利益分享机制设计研究》,《中国管理科学》第 1 期。
3. 董志强、严太华, 2007:《监察合谋: 惩罚、激励与合谋防范》,《管理工程学报》第 3 期。
4. 方阳春、姚先国, 2007:《高新企业薪酬制度研究》,《科研管理》第 5 期。
5. 蒋神州, 2011:《国有控股公司治理中合谋防御的机制设计》,《经济评论》第 1 期。
6. 李金波、聂辉华、沈吉, 2010:《团队生产、集体声誉和分享规则》,《经济学(季刊)》第 3 期。
7. 卢周来, 2009:《合作博弈框架下企业内部权力的分配》,《经济研究》第 12 期。
8. 钱雪松、孔东民, 2012:《内部人控制、国企分红机制安排和政府收入》,《经济评论》第 6 期。
9. Alchian, A., and H. Demsetz. 1972. "Production, Information Costs, and Economic Organization." *American Economic Review*, 62(5): 777-795.
10. Baiman, S., J. H. Evans III, and N. J. Nagarajan. 1991. "Collusion in Auditing." *Journal of Accounting Research*, 29(1): 1-18.
11. Bertomeu, J. 2007. "Can Labor Markets Help Resolve Collusion." *Economics Letters*, 95(3): 355-361.
12. Boylan, R. T. 1998. "Coalition-proof Implementation." *Journal of Economic Theory*, 82(1): 132-143.

13. Brown, M. , and E. Wolfstetter. 1989. "Tripartite Income – Employment Contracts and Coalition Incentive Compatibility." *The RAND Journal of Economics*, 20(3) :291 – 307.
14. Che, Y. K. ,and J. Kim. 2009. "Optimal Collusion – proof Auctions." *Journal of Economic Theory*, 144(2) :565 – 603.
15. Delgado, J. , and D. Moreno. 2004. "Coalition – proof Supply Function Equilibria in Oligopoly." *Journal of Economic Theory*, 114(2) :231 – 254.
16. Eswaran, M. , and A. Kotwal. 1984. "The Moral Hazard of Budget – breaking." *RAND Journal of Economics*, 15(4) :578 – 581.
17. Forges, F. , J. F. Mertens, and R. Vohra. 2002. "The Ex Ante Incentive Compatible Core in the Absence of Wealth Effects." *Econometrica*, 70(5) ,1865 – 1892.
18. Green, J. , and J. J. Laffont. 1979. "On Coalition Incentive Compatibility." *Review of Economic Studies*, 46 (2) : 243 – 254.
19. Grimaud, F. , J. J. Laffont, and D. Martimort. 2003. "Collusion, Delegation and Supervision with Soft Information." *Review of Economic Studies*, 70(2) :253 – 279.
20. Hart, S. , and M. Kurz. 1983. "Endogenous Formation of Coalitions." *Econometrica*, 51(4) :1047 – 1064.
21. Holmstrom, B. 1982. "Moral Hazard in Teams." *Bell Journal of Economics*, 13(2) :324 – 340.
22. Jackson, M. O. , and H. F. Sonnenschein. 2007. "Overcoming Incentive Constraints by Linking Decisions." *Econometrica*, 75(1) :241 – 257.
23. Kamishiro, Y. 2011. "Informational Size and the Incentive Compatible Coarse Core in Quasilinear Economies." *Games and Economic Behavior*, 71(2) :513 – 520.
24. Laffont, J. J. , and D. Martimort. 1997. "Collusion under Asymmetric Information." *Econometrica*, 65(4) :875 – 911.
25. Mas – Collé, A. , Whinston, M. D. , and J. R. Green. 2005. *Microeconomic Theory*, 673 – 684. Shanghai: Shanghai University of Finance & Economics Press.
26. Maskin, E. 2011. "Commentary: Nash Equilibrium and Mechanism Design." *Games and Economic Behavior*, 71(1) :9 – 11.
27. Moreno, D. , and J. Wooders. 1996. "Coalition – Proof Equilibrium." *Games and Economic Behavior*, 17(1) :80 – 112.
28. Sengupta, A. , and K. Sengupta. 1996. "A Property of the Core." *Games and Economic Behavior*, 12(2) :266 – 273.

Coalition – Proof Incentive Mechanism in Team Production

Dai Jiansheng

(Faculty of Management and Economics, Kunming University of Science and Technology)

Abstract: By means of the basic idea of “core”, this research establishes a coalition – proof mutually incentive model by introducing coalition – proof incentive compatibility constraints, and this paper characterizes the first – order optimal conditions and domain of its solutions. An incentive mechanism is coalition – proof only when every team member can acquire more in the grand coalition than in all possible sub – coalitions involved. Furthermore, it provides some interpretations, different from classic theories, for some economic phenomena such as salary structure, efficient wage, and insider control and so on. The economic reasons behind the phenomenon are partly attributed to bargaining powers and coalition likelihoods of some members. It points out that more attention should be paid to some work teams with coalition likelihood in mechanism designing. It can be applied to instruct teams with more than two members to design mechanism to incentive members’ efforts and prevent some members from forming an alliance.

Key Words: Team Production; Coalition – proof; Incentive Mechanism; Core

JEL Classification: C78, D74, D82

(责任编辑:赵锐、彭爽)