

Incentivos (Parte I)

Microeconomía II

FCEA - FCS

5 de mayo de 2026

Maestría en Economía (Opción Análisis Económico)



Ciencias Sociales
Universidad de la República
URUGUAY

Objetivos

1. Incentivos

- 1.1 Determinar los problemas de información asimétrica
- 1.2 Entender cómo corregir estas fallas de mercado... y cómo convivir con ellas.

2. Hashtags: #eficiencia, #incentivos, #hacemosloquepodemos, #lavidaescomplicada...

Previo

- Antes que nada, vamos a ver un video que explica todo este tema:
- Susana Olaondo - Julieta, ¿qué plantaste?
- ¡¡ Este video está “basado” en el trabajo de Stiglitz (1974): “Incentives and Risk Sharing in Sharecropping” !!

Previo

- Antes que nada, vamos a ver un video que explica todo este tema:
- Susana Olaondo - Julieta, ¿qué plantaste?
- ¡¡ Este video está “basado” en el trabajo de Stiglitz (1974): “Incentives and Risk Sharing in Sharecropping” !!

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica?

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas?

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas? Riesgo Moral y selección adversa.

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas? Riesgo Moral y selección adversa.
 - ¿Qué pasa cuando hay múltiples tareas?

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas? Riesgo Moral y selección adversa.
 - ¿Qué pasa cuando hay múltiples tareas? Difícil incentivar a los agentes.

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas? Riesgo Moral y selección adversa.
 - ¿Qué pasa cuando hay múltiples tareas? Difícil incentivar a los agentes.
 - ¿Qué pasa en entornos dinámicos?

¿Qué queremos entender?

- Muchas cosas:
 - ¿Qué pasa cuando hay información asimétrica? (todo mal, (casi) chau primer óptimo... :-)
 - ¿Cuáles son los principales problemas? Riesgo Moral y selección adversa.
 - ¿Qué pasa cuando hay múltiples tareas? Difícil incentivar a los agentes.
 - ¿Qué pasa en entornos dinámicos? ¡El compromiso importa! (Renegociación, expropiación).

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada
Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo del Principal

Ejemplos

Extensiones

Extensión I: mas de una tarea

Extensión II: múltiples agentes

Incentivos

- Estructura general:
 - Agente (Principal) **delega** en otro (Agente) la realización de una tarea.
 - La relación entre las partes genera **valor**.
 - Sin embargo, ambos agentes tienen **objetivos** diferentes.
 - El **resultado** de la acción del Agente impacta sobre el bienestar del principal.

Incentivos

- Estructura general:
 - Agente (Principal) **delega** en otro (Agente) la realización de una tarea.
 - La relación entre las partes genera **valor**.
 - Sin embargo, ambos agentes tienen **objetivos** diferentes.
 - El **resultado** de la acción del Agente impacta sobre el bienestar del principal.
 - ... Agente tiene **información privada** sobre alguna variable (ej.: tipo, esfuerzo).

Quien mueve / contrato

- Hay dos mecanismos según quien pueda mover en el juego.
 - Si mueve la parte informada: **señalar** (signalling).
 - Si mueve la parte no informada: **cernir** (screening).
- En riesgo moral, el problema es **post**-contractual: firmado el contrato, el Agente cambia su comportamiento (ej. no trabaja).
- En selección adversa, el problema es **pre**-contractual: los Agentes son lo que son, quiero que revelen sus características o las señalen.

Mas variantes...

- Un Principal, un Agente, estático, unidimensional. Problema: renta por información, riesgo.
- Un Principal, un Agente, estático, multidimensional. Problema: incentivos a realizar varias tareas o producir varios bienes.
- Un Principal, un Agente, dinámico. Problema: compromiso y renegociación.
- Un Principal, varios Agentes (grupos). Problemas: free-riding entre agentes, competencia entre agentes, colusión.
- Varios Principales (agencia común), un Agente. Problemas: free-riding entre principales, falta de coordinación, monitoreo inconsistente.

Nota...

- El texto que les dejé no presenta gran parte de lo que se ve acá :’(
- Usé el texto de Bolton y Dewatripont (2005): Contract Theory, que es muy avanzado.
- (Intenté) Expuse los modelos de la forma más sencilla posible.
- Donde no fue posible, sólo les dejé las ideas del modelo y sus conclusiones.
- También tienen el libro de Kreps (2018): The Motivation Toolkit, que es excelente.
- La idea, es invitarlos a pensar... el mensaje es: **no hay soluciones fáciles** (como la vida misma).

Digresión: incentivos



Digresión: incentivos

- Las motivaciones pueden ser: intrínsecas, o extrínsecas.
- **Intrínsecas:** hago algo porque es interesante o agradable en sí mismo.
 - Relevantes para los trabajadores del sector público (sentido de misión, características pro-sociales de los individuos, autonomía, carrera).
- **Extrínsecas:** al hacer algo puedo obtener objetivos distintos a lo que hago, como salario o reconocimiento.
 - Ojo: ¡no sólo el salario es relevante!

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Idea

- Supongamos un Principal neutral al riesgo que delega en un Agente para que realice una tarea (en general, una acción a). (¿Porqué delegaría, Leandro???)
- El resultado depende del esfuerzo (e) del Agente.
- El esfuerzo no es observable por el Principal, y está correlacionado con el resultado.
- No se puede contratar esfuerzo, pero si producto (y).
- Hay dos resultados (1, por bueno o alto; 0, por malo o bajo).
- **Problema:** balance entre riesgos e incentivos.

Una breve reflexión....

- Si el esfuerzo es observable, los problemas se resuelven con **supervisión**.
- Sin embargo, en muchos casos eso no es posible: el Principal no puede monitorear el esfuerzo.
- Lo que vamos a ver es cómo **alinear los incentivos del Agente con los del Principal**, cuando supervisar es costoso para el Principal.
- Ejemplos: cumplir horario no dice nada sobre esfuerzo (charlo con compañeros), estar sentado con los materiales de estudio no dice nada sobre aprendizaje (estoy distraído), estar en la computadora no dice nada sobre trabajo (estoy en Insta)... (sigo???)

Modelo: datos

1. Resultados: $y_1 > y_0$.
2. Agente: esfuerzo $e = \{0, 1\}$, costo de esfuerzo $c = \{0, c\}$, pago s (salario).

	e	$p(y = 1)$	$p(y = 0)$
3. Relación esfuerzo - éxito:	0 (bajo)	p_0	$1 - p_0$
	1 (alto)	p_1	$1 - p_1$

con $p_1 > p_0$.

4. Utilidad del principal: $U^P = \mathbb{E}(y) - s$ (neutral al riesgo).
5. Utilidad del agente: $U^A = s - c(e)$ (si neutral al riesgo), o $U^A = \sqrt{s} - c(e)$ (si averso al riesgo), utilidad de reserva u .

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Agente

- Supongamos que el esfuerzo es observable $\implies$ el Principal contrata “**esfuerzo**”.
- Supuesto: agente **neutral** al riesgo: $U^A = s$.

Si Principal quiere $e = 1$

- Principal maximiza

$$\max_{s_1} U_1^P \quad \max p_1 y_1 + (1 - p_1) y_0 - s_1$$

$$s.a \quad s_1 - c \geq u \quad (RP_1)$$

$$\Rightarrow s_1 = u + c.$$

- Principal obtiene: $U_1^P = p_1 y_1 + (1 - p_1) y_0 - (u + c).$

Si Principal quiere $e = 0$

- Principal maximiza

$$\max_{s_0} U_0^P \quad \max p_0 y_1 + (1 - p_0) y_0 - s_0$$

$$s.a \quad s_0 \geq u \quad (RP_0)$$

$\Rightarrow s_0 = u.$

- Principal obtiene: $U_0^P = p_0 y_1 + (1 - p_0) y_0 - u.$

¿Qué decide?

- Principal elige $e = 1$ si la utilidad esperada es mayor.

$$\Rightarrow U_1^P \geq U_0^P \iff p_1 y_1 + (1 - p_1) y_0 - (u + c) \\ \geq p_0 y_1 + (1 - p_0) y_0 - u \text{ y reordenando obtenemos:}$$

$$\underbrace{(p_1 - p_0)(y_1 - y_0) = \Delta p \times \Delta y}_{IM_{ge=1}} \geq \underbrace{c}_{CM_{ge=1}}$$

- Principal induce esfuerzo si $IM_{ge=1} \geq CM_{ge=1}$ (es decir, no siempre conviene inducir esfuerzo).
- El resultado no cambia si el Agente es averso al riesgo (cambian $s = \sqrt{s}$).

Resultado

- Si no hay asimetría de información:
 1. Principal asegura al Agente: le paga salario según esfuerzo, **independientemente del riesgo del resultado**.
 2. El Agente recibe un salario que paga su costo de oportunidad y, en su caso, de esforzarse.
 3. Principal induce esfuerzo sólo si es óptimo (si, a veces es eficiente “que la gente no haga nada”).
 4. El Principal corre todo el riesgo y se queda con el resultado residual (*residual claimant*).

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Agente

- Supongamos que el esfuerzo es **no** observable $\implies$ el Principal contrata “resultado” y .
 - Supuesto: agente **neutral** al riesgo: $U^A = s$.
 - **Problema:** el resultado anterior no implementa esfuerzo.
 - Si el Principal quiere implementar esfuerzo $e = 1$, como el esfuerzo no es observable, el Agente “dice” que se esfuerza pero no le conviene hacerlo.
 - Si se esfuerza $U(s_1|e = 1) = s_1 - c = u$,
 - Si no se esfuerza $U(s_1|e = 0) = s_1 - 0 = u + c > u$
- $\Rightarrow$ al Agente le conviene **no** esforzarse.

Información asimétrica

- El resultado de no esforzarse no cambia: le pago la utilidad de reserva $s_0 = u$.
- Para que se esfuerce, hay que pasar riesgo al agente $\implies$ tengo que hacer el salario **contingente** al resultado y .
- Además, para el Agente debe ser preferible esforzarse $\implies$ necesito una nueva restricción que haga $U_1^A > U_0^A$: es la **restricción de compatibilidad de incentivos**.

Compatibilidad de incentivos

Definición

La **compatibilidad de incentivos** garantiza que la utilidad del Agente asociada a una determinada acción a es mayor que cualquier acción alternativa a' .

Es decir $U(a) > U(a')$.

Principal induce $e = 1$

- Principal maximiza

$$\max_{s_1^1, s_0^1} p_1 (y_1 - s_1^1) + (1 - p_1) (y_0 - s_0^1)$$

$$s.a \quad p_1 s_1^1 + (1 - p_1) s_0^1 - c \geq u \quad (RP_1)$$

$$s.a \quad p_1 s_1^1 + (1 - p_1) s_0^1 - c - u \geq p_0 s_1^1 + (1 - p_0) s_0^1 - u \quad (RCI_1)$$

⇒ despejando en (RP_1) y (RCI_1) obtenemos:

$$s_1^1 = u + c \frac{(1-p_0)}{(p_1-p_0)} \text{ y } s_0^1 = u - c \frac{p_0}{(p_1-p_0)}.$$

- Principal obtiene: $U_1^P = p_1 y_1 + (1 - p_1) y_0 - (u + c)$, ¡lo mismo que antes!

Principal induce $e = 0$

- Principal maximiza

$$\max_{s_0} p_0 y_1 + (1 - p_0) y_0 - s_0$$

$$s.t. \quad s_0 \geq u \quad (RP_0)$$

$$\Rightarrow s_0 = u.$$

- Principal obtiene: $U_0^P = p_0 y_1 + (1 - p_0) y_0 - u.$

Comparación

- El problema con y sin información asimétrica es similar (si Principal y Agente neutral al riesgo):
 - el beneficio del Principal es igual en ambos casos $\implies$ la decisión de implementar esfuerzo no cambia por la información asimétrica.
- Cuando hay información asimétrica, el Principal tiene que **pasar riesgo** para inducir esfuerzo.
- El pago del Agente cuando se esfuerza, es un salario base (u) y un pago por resultados $\left\{ -c \frac{p_0}{(p_1 - p_0)}, +c \frac{(1 - p_0)}{(p_1 - p_0)} \right\}$.
- Esto **cambia** cuando el Agente es averso al riesgo.

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

**Nota I: Agente tiene
responsabilidad limitada**

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Restricción para el Agente

- La solución para implementar salarios implica una penalidad si $y = y_0$, $s_0^1 = u - c \frac{p_0}{(p_1 - p_0)}$.
- Supongamos que $s_0^1 \geq 0$, es decir el Agente tiene responsabilidad limitada.
- Si se cumple que $u < c \frac{p_0}{(p_1 - p_0)}$, entonces $s_0^1 = 0 \implies$ como $s_1^1 = s_0^1 + \frac{c}{(p_1 - p_0)} \implies s_1^1 = \frac{c}{(p_1 - p_0)}$.
- Estos valores inducen esfuerzo al Agente ... siempre que el Principal quiera.

Principal

- Antes, el Principal si quería inducir esfuerzo, obtenía:
$$U_1^P = p_1 y_1 + (1 - p_1) y_0 - (u + c).$$
- Ahora, obtiene:
$$U_1^{PrI} = p_1 \left(y_1 - \frac{c}{(p_1 - p_0)} \right) + (1 - p_1) y_0$$
$$= p_1 y_1 + (1 - p_1) y_0 - \frac{p_1 c}{(p_1 - p_0)}.$$
- Noten que si $u < c \frac{p_0}{(p_1 - p_0)} \implies U_1^{PrI} < U_1^P.$

Solución

- Si el Agente tiene responsabilidad limitada, el Principal no puede castigarlo tanto en el estado malo de la naturaleza.
 - Para que la *RCI* se siga cumpliendo, si sube el pago en el mal estado, tiene que bajarlo en el bueno.
 - Ello es costoso para el Principal, que obtiene un beneficio menor.
- ⇒ El Principal elige inducir esfuerzo ($e = 1$) menos veces.

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Solución

- Supongamos que el esfuerzo es **no** observable $\implies$ el Principal contrata “resultado” y.
- Supuesto: agente **averso** al riesgo: $U^A = \sqrt{s}$.
- Noten que si llamamos a $\sqrt{s} = S$, el análisis anterior se mantiene $\implies$ para encontrar la solución hay que hacer raíz.

$$\Rightarrow s_1^1 = \left(u + c \frac{(1-p_0)}{(p_1-p_0)}\right)^2 \text{ y } s_0^1 = \left(u - c \frac{p_0}{(p_1-p_0)}\right)^2.$$

Aversión al riesgo y Primer óptimo

- Con aversión al riesgo el Principal tiene que pagar **salarios mayores** que cuando el Agente es neutral al riesgo.
- ⇒ la **aversión al riesgo hace más difícil implementar el primer óptimo.**
- Se puede demostrar que (pero creanme!):

$$U^P(ANR) > U^P(AAR)$$

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Relación entre esfuerzo y resultados

- Basado en Holmstrom (1979): "Moral Hazard and Observability".
- Para que incentivar esfuerzo funcione, tiene que haber una relación "razonable" entre el esfuerzo y el resultado.
 - **Señal**: la parte del resultado asociada al esfuerzo del Agente.
 - **Ruido**: la parte del resultado que **no** depende del Agente.
- En el modelo, la señal es $\Delta p = (p_1 - p_0)$: cuanto aumenta el resultado positivo debido al esfuerzo.
- Sin embargo, aún con esfuerzo, no está garantizado el éxito: aún con esfuerzo, hay "ruido" $Var(p_1) = p_1(1 - p_1)$.

Digresión: SNR

- Se puede definir el “Cociente Ruido Señal” (Signal-to-Noise Ratio, SNR):

$$SNR = \frac{\mu^2}{\sigma^2} = \frac{(p_1 - p_0)^2}{p_1(1 - p_1)}$$

- Numerador: cuanto aumenta el resultado bueno por el esfuerzo.
- Denominador: la varianza del esfuerzo alto (ruido).

¿Y?

- Cuanto mayor $(p_1 - p_0)$, mayor la información que aporta y_1 sobre e_1 .
- Si $p_1 \approx p_0 \implies$ el esfuerzo aporta poco: el resultado es todo “ruido”.
- Si $p_1 \gg p_0 \implies$ el esfuerzo aporta mucho al resultado: todo señal.

¿Qué es “ruido”?

- El ruido son aquellas cosas **fuera de control** del Agente y que influyen en el resultado.
- Supongamos que el resultado es “estudiante pasa de año”, ¿qué influye?
 - Si los padres lo manden a la escuela,
 - Si estudian con él,
 - Si tiene un lugar para estudiar y materiales,
 - Si la escuela está cerca, y el lugar no es riesgoso,
 - Si come todos los días,
 -
- En estos casos, es poco probable que el resultado docente esté correlacionado con el resultado: hay mucho “ruido”

El ruido “molesta”

- Vean que si $p_1 \approx p_0 \implies p_1 \rightarrow p_0$: el resultado varía poco con el esfuerzo.

- Volvamos a los salarios: $s_1^1 = u + c \frac{(1-p_0)}{(p_1-p_0)}$ y

$$s_0^1 = u - c \frac{p_0}{(p_1-p_0)}:$$

- si $p_1 \rightarrow p_0 \implies (p_1 - p_0) \rightarrow 0 \implies c \frac{(1-p_0)}{(p_1-p_0)} \rightarrow \infty!$

$\Rightarrow$ se cumple que $s_1^1 \rightarrow \infty$ y $s_0^1 \rightarrow -\infty$.

- Esto es costoso para el Principal.

No incentivar

- Si la señal es poco informativa, hay que dar incentivos (y penalizaciones) “desmedidos”.
- ⇒ **no conviene incentivar esfuerzo.**
- Por esto es **MUY importante** tener claro el “ruido”: lo que afecta el resultado y está fuera del control del Agente.
 - Si no podemos mitigar estos efectos fuera de control ... los incentivos no ayudan.

Función objetivo

- Basado en Baker (1992): "Incentive Contracts and Performance Measurement".
- Cuando hay información asimétrica sobre el esfuerzo, el Principal contrata sobre el producto y (el resultado).
 - Principal $U^P(y, e)$
- ¿Qué pasa cuando no se puede contratar el resultado?
 - Principal $U^P(e, \varepsilon)$, con ε el estado de la naturaleza (shocks).
- Esto pasa cuando el objetivo del principal no es observable y verificable por terceros:
 - Ej. ¿cuál es el valor de una política pública?, ¿cuál es el valor futuro de una patente?, ¿cuál es el margen de las ventas?

Contrato

- Cuando no se puede contratar el resultado, se contrata una medida de desempeño alternativa $P(e, \varepsilon)$.
- Sin embargo, al ¡Principal le interesa $U^P(e, \varepsilon)$!
- Cuanta mayor correlación entre P y el resultado sobre U^P , mejor será el contrato.
- Es decir, hay una distancia entre la métrica P y el objetivo del Principal U^P .

Ejemplos

- Ejemplo I:
 - Principal: gobierno; Objetivo: “Educación de calidad”.
 - Métrica: cantidad de estudiantes que terminan educación primaria.
 - Incentivo: no tomar pruebas, que los estudiantes pasen sin estudiar.
- Ejemplo II:
 - Principal: empresa; Objetivo: “Patentar productos”.
 - Métrica: cantidad de patentes de la empresa.
 - Incentivo: esforzarse en patentar muchos procesos sencillos, en vez de productos relevantes.



Problemas

- Cuando el objetivo del Principal no puede contratarse, es difícil alinear los incentivos del Agente.
- Agente tiene incentivos a (recordar que esfuerzo es no observable):
 - (Gaming) manipular las medidas de desempeño (ej. patentar cosas fáciles vs. valiosas).
 - (Descremar) elegir solo casos fáciles para mejorar resultados (ej. dar clase a sólo a los alumnos buenos).
- (Si les interesa este tema, les recomiendo un libro muy bueno: Jerry Muller "The Tyranny of Metrics").



Dueño tierra y arrendador (aparcero)

- Este es el paper clásico de Stiglitz, 1974.
- Principal: dueño de la tierra :(.
- Agente: campesino que la cultiva.
- Problema: el campesino resuelve e, y el dueño de la tierra lo observa parcialmente.
- Solución: pasar riesgo, pago en función de la cosecha (s_1^1, s_0^1) .

Gobierno y contratista

- Principal: Estado.
- Agente: contratista que tiene que llevar a cabo una obra.
- Problema: el contratista resuelve a , que en este caso puede ser usar materiales baratos o de mala calidad.
- Solución: no hay forma de pasar riesgo, hay que atar pagos a cumplir etapas y verificar calidad.



Teoría de la empresa

- Este es el trabajo clásico de Jensen y Meckling (1976): "Theory of the firm: Managerial behavior, agency costs and ownership structure".
- Principal: Accionista externo (inversor).
- Agente: Gerente, quien también puede ser propietario parcial de la empresa.
- Problema: el Gerente puede desviar recursos de la empresa para su beneficio privado.
- Solución: aumentar la participación del Gerente en la empresa (hacerlo accionista).

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Idea

- Vamos a analizar los problemas de **múltiples tareas** (basado en Holmstrom y Milgrom (1991): "Multitask Principal-Agent Analyses: Incentive Contracts, Asset Ownership, and Job Design").
- En general, los Agentes tienen que realizar más de una tarea.
- Si hay varias tareas $\implies$ el problema no es solo de incentivos versus riesgo: importa como los incentivos a realizar una tarea afectan a las restantes.

Ejemplo

- Atar la remuneración de los docentes a los resultados de los estudiantes en los test.
 - test proveen una medida independiente de desempeño :)
 - no miden otras tareas docentes (ej. escritura, habilidad para discutir, para filosofar, para socializar) :(
- ⇒ atar remuneración hace que los docentes se enfoquen **sólo** (o principalmente) en los resultados de los test.

Idea (cont.)

- Noten que este problema es similar al de riesgo moral ya visto.
- El agente tiene una distribución de probabilidades sobre resultados y .
- ¡El Principal elige que incentivo da y eso determinar el resultado!

Estructura: supuestos

- El modelo, para que tenga solución, es bastante complicado.
- Supuestos:
 - dos tareas y_1, y_2 , cada una requiere su esfuerzo e_1, e_2 ,
 - cada tarea recibe shocks $\varepsilon_i \sim N(0, \Sigma)$, y los shocks entre tareas están correlacionados $\rho \in [-1, 1]$,
 - Agente es averso al riesgo (utilidad CARA
 $u(x) = -e^{-\gamma w}$),
 - el pago es $w = a + s_1 y_1 + s_2 y_2$ (fijo, mas variables).

Estructura: sustitución de esfuerzos

- el Agente tiene costos cuadráticos en ambas tareas $\frac{1}{2}c(e_1^2 + e_2^2) + \delta e_1 e_2$, donde δ mide el grado de sustitución entre tareas:
 - Si $\delta = 0$, las tareas son independientes (el problema es similar a lo visto antes).
 - Si $\delta > 0$, hay un problema de **sustitución de esfuerzos**: esforzarse en una tarea sube el CMg de hacer la otra tarea.

Resumen de resultados

- En equilibrio, $e_i = \frac{s_i c - \delta s_j}{c^2 - \delta^2}$.
 1. Si $\delta = 0$, el Agente se esfuerza en tarea i siempre que el pago $s_i > 0$.
 2. Si $\delta > 0$, hay sustitución de esfuerzos: cualquier pago en la tarea j disminuye el esfuerzo en la tarea i .

Salarios complementarios

- **Salarios** son una función de $(c, \delta, \gamma, \sigma^2)$ y **son complementarios entre sí** (debido a la sustitución de esfuerzos).

Salarios complementarios: ejemplos

1. Si la tarea 2 no está controlada por el Principal (trabajo fuera de la empresa) y $\uparrow s_2 \implies$ el Principal tiene que $\uparrow s_1$ para que el Agente se esfuerce en la tarea 1.
2. Si al Principal le interesa la tarea 1, pero y_1 es más difícil de medir para el Principal ($\uparrow \sigma_1^2$) $\implies$ tengo que $\downarrow s_1$ ¡pero también $\downarrow s_2$! (sino el Agente se pasa a la tarea 2, que no interesa).

En este caso, conviene pagar el fijo y no incentivar ningún esfuerzo (incentivos de bajo poder).

Estructura: conflicto entre tareas

- Una alternativa es que llevar adelante una tarea **dificulte obtener el otro resultado**: $P(q_i = 1) = \alpha + \rho e_i - \gamma e_j$.
- En este caso, esforzarse en la tarea j disminuye la probabilidad de éxito de la tarea i .
- Escenarios:
 1. Contratar un agente que realice ambas tareas $\implies$ salario extra si **ambas** tareas son exitosas.
 2. Contratar dos agentes, que realicen cada uno una tarea (noten que lo que hace el otro dificulta lo que hago yo!) $\implies$ salario extra si uno es exitoso.

Defensores

- Si γ es bajo, contrato un agente; si alto, contrato dos (ej: defensores, *advocates*).
- Típico ejemplo: la justicia, hay un abogado defensor y un fiscal que acusa (dos agentes que manejan información opuesta sobre una tarea).

Índice

Presentación

Riesgo moral

Primer óptimo

Segundo óptimo: Agente
neutral al riesgo

Nota I: Agente tiene
responsabilidad limitada

Segundo óptimo: Agente
averso al riesgo

Nota II: señal y ruido

Nota III: La función objetivo
del Principal

Ejemplos

Extensiones

Extensión I: mas de una
tarea

Extensión II: múltiples
agentes

Idea

- Vamos a analizar los problemas de los **trabajos en equipo** (basado en Holmstrom, 1982).
- Supongamos **dos** Agentes (A, B) que ejercer esfuerzo (e) para obtener **un** producto $y = y(e^A, e^B)$.
- El resultado depende que **ambos** Agentes se esfuercen.
- Los Agentes forman una **sociedad**, es decir se reparten el producto en partes iguales: $w^A = w^B = \frac{y}{2}$.
- La producción (y) que tiene dos resultados (y_1 , por bueno o alto; $y_0 = 0$, por malo o bajo).

Idea (cont.)

- Relación esfuerzo - éxito:

$$P(y = y_1 | e_A, e_B) = \begin{cases} 0, & (0,0) \\ q, & (1,0) \text{ o } (0,1) , \\ p, & (1,1) \end{cases} \quad p > q > 0$$

Solución conjunta

- Si los agentes pudieran coordinar, entonces maximizarían

$$\max_{e_A, e_B} E(U_A^A + U_B^A)$$

- Eligen $e_A = e_B = 1 \iff \underbrace{py_1 - 2c}_{\text{ben. esf. ambos}} > \underbrace{qy_1 - c}_{\text{ben. esf. uno}}.$
- Es decir, si $\underbrace{(p - q)y_1}_{\text{IMg}_{(0,1) \rightarrow (1,1)}^{\text{conj}}} > \underbrace{c}_{\text{CMg}_{(0,1) \rightarrow (1,1)}}.$

Solución individual

- Si los agentes no pueden coordinar, entonces maximizan:

$$\max_{e_i} E(U_i^A) \quad i = A, B$$

- Eligen $e_i = 1 | e_j = 1 \iff \underbrace{(p - q) \frac{y_1}{2}}_{IMg_{(0,1) \rightarrow (1,1)}^{ind}} > \underbrace{c}_{CMg_{(0,1) \rightarrow (1,1)}} .$

- Noten que $IMg^{conj} = (p - q) y_1 > IMg^{ind} = (p - q) \frac{y_1}{2} !$

Conclusiones

- Los Agentes tienen **menos** incentivos a esforzarse.
- Trabajar en un equipo es similar a colaborar en un **bien público**: no me apropio de todo el beneficio.
- El problema se resuelve si un tercero (Principal) distribuye las tareas y retiene el **beneficio residual**.
- La clave es que las participaciones de los Agentes no sumen uno (*budget breaking*).
- Vean Ejercicio sobre el tema.

Alternativas: competencia

- Una forma de incentivar esfuerzo cuando hay múltiples agentes es realizar torneos (*tournaments*).
- Para ello se requiere que **el esfuerzo de cada Agente sea observable** (es decir, no hay un producto conjunto).
- Un torneo determina un pago relativo al que tiene el mayor producto q_i .
- Es utilizado habitualmente para establecer ascensos, cartas de recomendación, ... y en todos los deportes.
- Problema: los torneos generan incentivos perversos, por ejemplo sabotaje.

Alternativas: ayuda

- Una alternativa a la competencia es la ayuda.
- De nuevo, el producto de cada Agente es observable.
- Los Agentes pueden destinar parte de su esfuerzo (e) a **ayudar a otros**.
- Como el esfuerzo es costoso para el Agente $\implies$ el Principal tiene que pagar al Agente A si al Agente B le va bien (esfuerzo no es observable!).

Alternativas: ayuda (cont.)

- Problema: dificulta la especialización (hay que ayudar al otro...).
- Si no hay costos de especialización $\implies$ siempre tiene sentido ayudar.
- Si la especialización es muy beneficiosa $\implies$ tiene sentido ayudar $\iff$ hace mucha diferencia en el resultado.
 - En equilibrio hay mucha cooperación entre los Agentes, o ninguna.

Alternativas: cooperación/colusión

- Si los Agentes pueden observar el esfuerzo, pero el Principal no $\implies$ tiene sentido que coordinen entre sí.
- Funciona si el producto de cada Agente no está sometido a un shock común (si sus resultados no están muy correlacionados).
- En estos casos, los Agentes se **controlan** unos a otros.
- Si los resultados están muy correlacionados, entonces conviene que **compitan**.

Hasta aquí por ahora

Próximo: Incentivos Parte II