

Salir a la pizarra: Un modelo de monitoreo docente bajo IA generativa

Joaquín Martínez

Facultad de Economía y Negocios, Universidad de Chile

Otoño 2026

La IA Agéntica bien utilizada permite producir entregables académicos a un costo casi nulo.

Tarea entregada $\neq$ Conocimiento adquirido

La respuesta natural del cuerpo docente ha sido más exigencias en forma de: presentaciones aleatorias, exámenes orales, etc.

Sin embargo esta respuesta ha sido intuitiva, no necesariamente fundamentada **¿cuánto monitorear? ¿cómo diseñarlo? ¿qué parámetros importan?**

¿Cuál es el nivel óptimo de monitoreo para un profesor que maximiza el aprendizaje de sus estudiantes?

Modelamos el monitoreo en forma de presentaciones aleatorias en clase.

Trade-off:

- El monitoreo induce esfuerzo genuino y disuade la delegación a la IA
- Pero consume tiempo valioso de clase

- Cotton et al. (2024) fueron de los primeros quienes se preocuparon de ChatGPT-3. Plantearon mecanismos que, para las capacidades de la IA hoy en día **quedaron obsoletos**.
- Si bien ahora se reconoce como herramienta se advierten sus posibles consecuencias en **depreciación de habilidades (Huang & Vishnoi, 2026; Huang et al., 2026) y efectos dañinos en aprendizaje** (Bastani et al., 2025).

Ninguno modela la elección óptima de monitoreo del docente como respuesta estratégica, lo que da pie a nuestro aporte.

Estructura del juego base

Modelamos un juego secuencial con información perfecta y completa de un período (tipo Stackelberg).

- **Jugadores:** $\mathcal{N} = \{\text{Profesor, Estudiante}\}$
- **Orden del juego:** En la etapa $k = 1$ el Profesor diseña el mecanismo de monitoreo. En $k = 2$ el Estudiante observa el mecanismo y elige qué tanto esforzarse y cuánto delegar a la IA
- **Estrategias:** El Profesor elige $m \in [0, 1]$ interpretable como probabilidad de ser llamado a presentar. El Estudiante elige $e(m)$ y $a(m)$ no negativos como esfuerzo y delegación respectivamente.

El juego tiene etapas y por tanto buscamos el **equilibrio perfecto en subjuegos** por medio de inducción hacia atrás.

Pagos jugadores

La nota es la suma de e y a , la cual se valora bajo un parámetro $\alpha > 0$.

$$U^S(e, a \mid m) = \underbrace{\alpha(e + a)}_{\text{valoración nota}} - \underbrace{m\phi a}_{\text{costo exposición}} - \underbrace{\frac{e^2}{2(1 + \delta m)}}_{\text{costo esfuerzo}} - \underbrace{\frac{\gamma}{2} a^2}_{\text{costo IA}}$$

El profesor valora de manera lineal el esfuerzo del estudiante y el monitoreo tiene costos convexos

$$U^P(m) = e^*(m) - \frac{\kappa}{2} m^2$$

Equilibrio Perfecto en Subjuegos

Etapla 2: el estudiante tiene las mejores respuestas ante m :

$$e^*(m) = \alpha(1 + \delta m), \quad a^*(m) = \frac{\alpha - m\phi}{\gamma}$$

Si el monitoreo es productivo se aumenta el esfuerzo (Canal Pedagógico) $\partial e^*/\partial m = \alpha\delta > 0$.
La delegación se castiga como un impuesto (Canal disuasivo) $\partial a^*/\partial m = -\phi/\gamma < 0$.

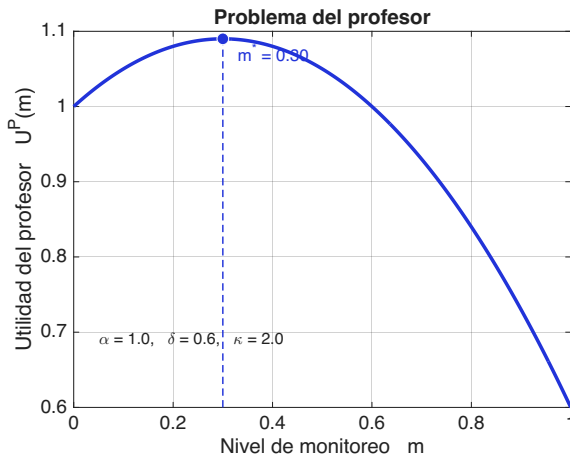
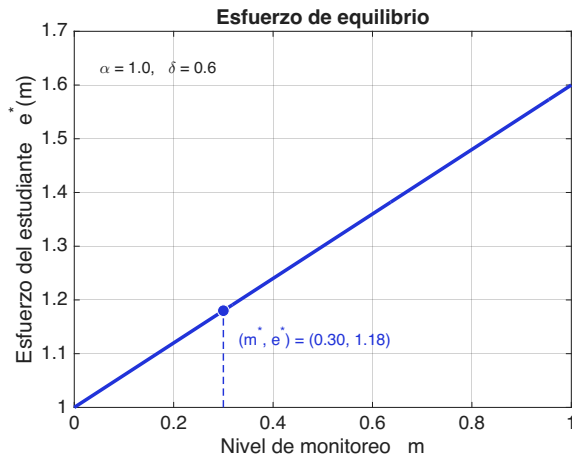
Etapla 1: profesor toma $e^*(m)$ y optimiza eligiendo

$$m^* = \frac{\alpha\delta}{\kappa}$$

El profesor iguala el beneficio marginal de monitorear (mayor conocimiento inducido) con el costo marginal (tiempo de clase perdido):

El equilibrio perfecto en subjuegos es $(m^*, e^*(m^*), a^*(m^*))$.

Equilibrio del modelo base



Al profesor le gustaría monitorear más para inducir más esfuerzo, pero el costo de oportunidad de clase empieza a ponderar por sobre el beneficio

Implicancias normativas del modelo base

El monitoreo óptimo no depende de la severidad del castigo ni del costo de la IA. La severidad (ϕ) disuade la delegación pero no motiva el estudio. Para aumentar el aprendizaje hay que

Hacer la presentación más productiva

($\uparrow \delta$):

- Exigir explicación del razonamiento, no lectura de la respuesta
- Preguntas conceptuales abiertas
- Pedir conexiones con clases previas
- Hacer del error un momento pedagógico para la audiencia

Bajar el costo de oportunidad ($\downarrow \kappa$):

- Integrar la presentación al contenido de la clase
- Presentaciones cortas y focalizadas
- Usarlas como repaso activo, no como evaluación
- Generar externalidades positivas hacia el resto del curso

Conclusión y Limitaciones

Conclusión: Desde un punto de vista de teoría de juegos podemos formalizar el mecanismo por el cual inducir un mayor aprendizaje en los alumnos. Dado este marco teórico podemos derivar propuestas de política que mejoren la educación.

Limitaciones

- Asumimos que al estudiante no le importa aprender y que la IA no tiene efectos positivos sobre el aprendizaje
- El estudiante es representativo, no consideramos heterogeneidad
- No hay un tope en la nota y no hay targeting. Posible extensión: usar Goette et al. (2004) como función de utilidad donde hay una aversión a la pérdida, con respecto a cierta nota.

ANEXO 1: Modelo Extendido

Modelo extendido

El anterior modelo castigaba la delegación de la IA como un impuesto plano.

- Un approach más realista es que se castigue el gap entre el conocimiento (e) y el trabajo que hizo la IA (a).

El castigo esperado, respuesta en esfuerzo y el monitoreo serán:

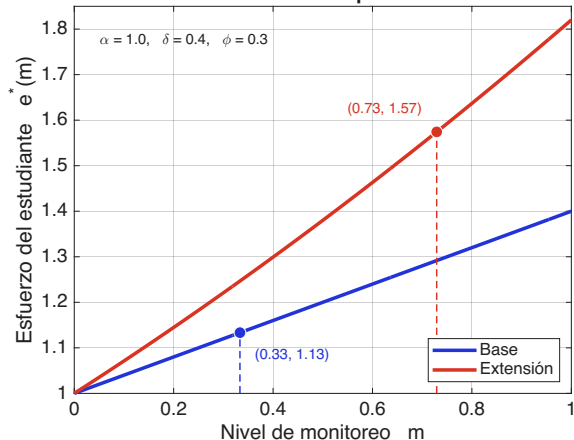
$$\alpha\phi m(a - e), \quad e^*(m) = \alpha(1 + \phi m)(1 + \delta m), \quad m^* = \frac{\alpha(\phi + \delta)}{\kappa - 2\alpha\phi\delta}.$$

Frente al modelo base $m^* = \alpha\delta/\kappa$, la severidad ahora entra por el numerador y el denominador vía la defensa de la nota.

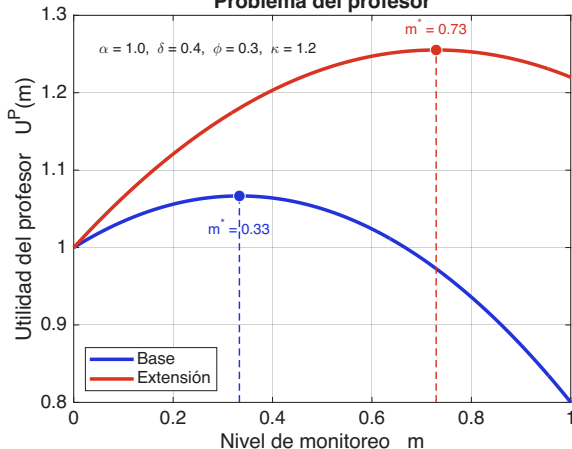
⇒ Mismas recomendaciones para δ y κ , pero se reconoce que la severidad del castigo si induce esfuerzo.

Mejores respuestas

Esfuerzo de equilibrio



Problema del profesor



Ahora el monitoreo óptimo es mayor porque el beneficio marginal en conocimiento (esfuerzo) aumentó dado el castigo sobre el gap.

Anexo 2: Derivaciones

Anexo: especificación formal del juego

Juego dinámico de información completa y perfecta en dos etapas (Stackelberg):

$$\Gamma = \langle \mathcal{N}, (A_i)_{i \in \mathcal{N}}, (U^i)_{i \in \mathcal{N}}, \text{orden temporal} \rangle$$

- **Jugadores:** $\mathcal{N} = \{P, S\}$, con $A_P = [0, 1]$ y $A_S = \mathbb{R}_+^2$
- **Pagos sobre perfiles de acciones:** $U^S(e, a \mid m)$ y

$$U^P(m, e, a) = f(e) - \frac{\kappa}{2}m^2, \quad f(e) = e$$

- **Estrategias:** $s_P = m \in [0, 1]$; $s_S = \sigma : [0, 1] \rightarrow \mathbb{R}_+^2$, $\sigma(m) = (e(m), a(m))$
- **Solución:** EPS por inducción hacia atrás — $\sigma^*(m)$ óptima en *todo* subjuego $m \in [0, 1]$; m^* óptimo anticipando σ^*

Anexo: concavidad del problema del Estudiante

El Hessiano de U^S respecto de (e, a) es

$$\nabla^2 U^S = \begin{pmatrix} -\frac{1}{1 + \delta m} & 0 \\ 0 & -\gamma \end{pmatrix}$$

Con $\delta > 0$, $\gamma \in (0, 1)$ y $m \in [0, 1]$ es definido negativo: U^S es estrictamente cóncava y las CPO caracterizan el máximo en todo subjuego.

En la extensión el término adicional $-\alpha\phi m(a - e)$ es lineal en (e, a) : el Hessiano no cambia.

Anexo: mejor respuesta completa (modelo base)

CPO del Estudiante, dado m :

$$\alpha - \frac{e}{1 + \delta m} = 0, \quad \alpha - m\phi - \gamma a = 0$$

La estrategia completa, definida para todo $m \in [0, 1]$:

$$e^*(m) = \alpha(1 + \delta m), \quad a^*(m) = \max\left\{0, \frac{\alpha - m\phi}{\gamma}\right\}$$

- $e^*(m) > 0$ siempre
- a^* interior $\iff m \leq \alpha/\phi$; si $m > \alpha/\phi$, esquina $a^*(m) = 0$

Anexo: etapa 1 y equilibrio (modelo base)

Forma reducida del Profesor:

$$U^P(m) = \alpha(1 + \delta m) - \frac{\kappa}{2}m^2, \quad \frac{d^2 U^P}{dm^2} = -\kappa < 0$$

CPO necesaria y suficiente: $\alpha\delta = \kappa m$. Si $\alpha\delta > \kappa$, esquina $m^* = 1$.

EPS único (interior, con $m^* \leq 1$ y $m^* \leq \alpha/\phi$):

$$m^* = \frac{\alpha\delta}{\kappa}, \quad e^*(m^*) = \alpha\left(1 + \frac{\alpha\delta^2}{\kappa}\right), \quad a^*(m^*) = \frac{\alpha}{\gamma}\left(1 - \frac{\delta\phi}{\kappa}\right)$$

Anexo: estática comparativa completa (modelo base)

$$\frac{\partial m^*}{\partial \alpha} = \frac{\delta}{\kappa} > 0, \quad \frac{\partial m^*}{\partial \delta} = \frac{\alpha}{\kappa} > 0, \quad \frac{\partial m^*}{\partial \kappa} = -\frac{\alpha \delta}{\kappa^2} < 0$$
$$\frac{\partial m^*}{\partial \phi} = \frac{\partial m^*}{\partial \gamma} = 0$$

- La severidad ϕ y el costo de la IA γ **no** afectan el monitoreo óptimo: solo entran en a^* , y el Profesor valora únicamente $f(e) = e$
- Sobre las mejores respuestas: $\partial e^* / \partial m = \alpha \delta > 0$ (canal pedagógico),
 $\partial a^* / \partial m = -\phi / \gamma < 0$ (canal disuasivo)

Anexo: extensión — problema del Estudiante

$$U^S(e, a \mid m) = \alpha(e + a) - \alpha\phi m(a - e) - \frac{e^2}{2(1 + \delta m)} - \frac{\gamma}{2}a^2, \quad \phi \in [0, 1]$$

CPO:

$$\alpha + \alpha\phi m - \frac{e}{1 + \delta m} = 0, \quad \alpha - \alpha\phi m - \gamma a = 0$$

Mejores respuestas interiores:

$$e^*(m) = \alpha(1 + \phi m)(1 + \delta m), \quad a^*(m) = \frac{\alpha(1 - \phi m)}{\gamma}$$

Ahora $\partial e^* / \partial \phi = \alpha m(1 + \delta m) \geq 0$: la severidad induce esfuerzo.

Anexo: extensión — problema del Profesor

Sustituyendo $e^*(m)$:

$$U^P(m) = \alpha + \alpha(\phi + \delta)m + \left(\alpha\phi\delta - \frac{\kappa}{2}\right)m^2, \quad \frac{d^2 U^P}{dm^2} = 2\alpha\phi\delta - \kappa$$

Bajo $\kappa > 2\alpha\phi\delta$, U^P es estrictamente cóncava y la CPO da

$$m^* = \frac{\alpha(\phi + \delta)}{\kappa - 2\alpha\phi\delta} > 0$$

Anexo: extensión — supuestos de regularidad

- ❶ **Concavidad:** $\kappa > 2\alpha\phi\delta$
- ❷ **Monitoreo interior:** $m^* \leq 1 \iff \kappa \geq \alpha(\phi + \delta + 2\phi\delta)$
- ❸ **Delegación no negativa:** $\phi m^* \leq 1$, garantizado por $\phi \in [0, 1]$ y $m^* \leq 1$
- ❹ **Gap no negativo:** $a^*(m^*) \geq e^*(m^*)$ (castigo, no subsidio)

Condición suficiente para 4: como $a^* - e^*$ es decreciente en m , basta evaluarla en $m = 1$:

$$\frac{1 - \phi}{\gamma} \geq (1 + \phi)(1 + \delta) \implies a^*(m) \geq e^*(m) \quad \forall m \in [0, 1]$$

Anexo: estática comparativa de la extensión

En el equilibrio interior, con $D = \kappa - 2\alpha\phi\delta > 0$:

$$\frac{\partial m^*}{\partial \phi} = \frac{\alpha(\kappa + 2\alpha\delta^2)}{D^2} > 0, \quad \frac{\partial m^*}{\partial \delta} = \frac{\alpha(\kappa + 2\alpha\phi^2)}{D^2} > 0$$

Esbozo (regla del cociente, $N = \alpha(\phi + \delta)$): el numerador de $\partial m^*/\partial \phi$ es

$$\alpha D + 2\alpha\delta N = \alpha\kappa - 2\alpha^2\phi\delta + 2\alpha^2\phi\delta + 2\alpha^2\delta^2 = \alpha(\kappa + 2\alpha\delta^2)$$

Contraste con el base: $\partial m^*/\partial \phi$ pasa de 0 a estrictamente positivo — castigar el gap convierte la severidad en incentivo al esfuerzo.

Referencias I

- Bastani, H., Bastani, O., Sungu, A., & Mariman, R. (2025). Generative ai without guardrails can harm learning: Evidence from high school mathematics. *Proceedings of the National Academy of Sciences (PNAS)*, 122(26), e2422633122. <https://doi.org/10.1073/pnas.2422633122>
- Cotton, D. R. E., Cotton, P. A., & Shipway, J. R. (2024). Chatting and cheating: Ensuring academic integrity in the era of chatgpt. *Innovations in Education and Teaching International*, 61(2), 228–239. <https://doi.org/10.1080/14703297.2023.2190148>
- Goette, L., Huffman, D., & Fehr, E. (2004). Loss aversion and labour supply. *Journal of the European Economic Association*, 2(2/3), 216–228. Retrieved June 8, 2026, from <http://www.jstor.org/stable/40004898>
- Huang, L., & Vishnoi, N. (2026). The geometry of learning under ai delegation [Available at SSRN]. *SSRN Electronic Journal*. <https://doi.org/10.2139/ssrn.6454543>
- Huang, L., Xiao, W., & Vishnoi, N. (2026). Delegation and verification under ai. <https://doi.org/10.2139/ssrn.6463045>

Apéndice técnico

Formalización y resolución del juego de monitoreo docente

Joaquín Martínez

16 de junio de 2026

1. El juego

Definición 1 (El juego Γ). El modelo es un juego dinámico de información completa y perfecta en dos etapas, con estructura de Stackelberg,

$$\Gamma = \langle \mathcal{N}, (A_i)_{i \in \mathcal{N}}, (\succeq_i)_{i \in \mathcal{N}}, \text{orden temporal} \rangle,$$

definido por:

1. **Jugadores.** $\mathcal{N} = \{P, S\}$, donde P es el Profesor (líder) y S el Estudiante (seguidor).
2. **Orden temporal.** En la etapa $k = 1$ el jugador P elige una acción, en la etapa $k = 2$ el jugador S , habiendo observado la acción de P , elige la suya.
3. **Conjuntos de acciones.** $A_P = [0, 1]$ y $A_S = \mathbb{R}_+^2$. Una acción de P es $m \in [0, 1]$ (probabilidad de monitoreo), una acción de S es un par $(e, a) \in \mathbb{R}_+^2$ (esfuerzo propio y delegación a la IA).
4. **Información.** El juego es de información perfecta: S observa m antes de elegir (e, a) . Los parámetros $(\alpha, \phi, \delta, \gamma, \kappa)$ son conocimiento común.
5. **Pagos.** Dados por las funciones $U^S : \mathbb{R}_+^2 \times [0, 1] \rightarrow \mathbb{R}$ y $U^P : [0, 1] \times \mathbb{R}_+^2 \rightarrow \mathbb{R}$ definidas en las Secciones 2 y 3.

Definición 2 (Estrategias). Como el juego es de información perfecta y P se mueve primero:

- Una estrategia (pura) del Profesor es una acción $s_P = m \in [0, 1]$.
- Una estrategia (pura) del Estudiante es una función $s_S = \sigma : [0, 1] \rightarrow \mathbb{R}_+^2$ que a cada nivel de monitoreo observado le asigna una respuesta, $\sigma(m) = (e(m), a(m))$.

Definición 3 (Concepto de solución). La solución es un equilibrio perfecto en subjuegos (EPS), un perfil $(s_P^*, s_S^*) = (m^*, \sigma^*)$ tal que $\sigma^*(m)$ es óptima para S en todo subjuego indexado por $m \in [0, 1]$, y m^* es óptima para P anticipando σ^* . El EPS se obtiene por inducción hacia atrás: primero se resuelve $\sigma^*(\cdot)$ y luego m^* .

Supuesto 1 (Parámetros). $\alpha > 0, \phi > 0, \delta > 0, \gamma \in (0, 1), \kappa > 0$.

2. Modelo base

El pago del Estudiante es

$$U^S(e, a \mid m) = \alpha(e + a) - m\phi a - \frac{e^2}{2(1 + \delta m)} - \frac{\gamma}{2} a^2, \quad (1)$$

y el del Profesor, definido sobre el perfil de acciones,

$$U^P(m, e, a) = f(e) - \frac{\kappa}{2} m^2, \quad f(e) = e. \quad (2)$$

En la forma reducida, evaluado en la mejor respuesta del Estudiante, $U^P(m) = f(e^*(m)) - \frac{\kappa}{2} m^2$.

2.1. Etapa 2: problema del Estudiante

Para que la solución del estudiante sea un máximo la función de utilidad tiene que ser estrictamente cóncava para todo $m \in [0, 1]$. Para comprobar esto derivamos el Hessiano de (1) respecto de (e, a) :

$$\nabla^2 U^S = \begin{pmatrix} \frac{\partial^2 U^S}{\partial e^2} & \frac{\partial^2 U^S}{\partial e \partial a} \\ \frac{\partial^2 U^S}{\partial a \partial e} & \frac{\partial^2 U^S}{\partial a^2} \end{pmatrix} = \begin{pmatrix} -\frac{1}{1 + \delta m} & 0 \\ 0 & -\gamma \end{pmatrix}.$$

Dado el Supuesto 1 $\nabla^2 U^S$ es definido negativo y U^S es estrictamente cóncava.

Encontramos la mejor respuesta del estudiante despejando las variables de control de cada condición de primer orden.

$$\frac{\partial U^S}{\partial e} = \alpha - \frac{e}{1 + \delta m} = 0, \quad \frac{\partial U^S}{\partial a} = \alpha - m\phi - \gamma a = 0.$$

Despejando se obtiene (3). Como $\alpha, \delta, m \geq 0$, $e^*(m) = \alpha(1 + \delta m) > 0$. Para a , $a^*(m) \geq 0 \iff \alpha - m\phi \geq 0 \iff m \leq \alpha/\phi$, si $m > \alpha/\phi$ el óptimo es la solución de esquina $a^*(m) = 0$. Entonces, si la solución es interior entonces

$$e^*(m) = \alpha(1 + \delta m), \quad a^*(m) = \frac{\alpha - m\phi}{\gamma}. \quad (3)$$

La solución interior $a^*(m) \geq 0$ requiere $m \leq \alpha/\phi$ y $e^*(m) > 0$ siempre. En esta región el monitoreo eleva el esfuerzo propio y reduce la delegación

$$\frac{\partial e^*}{\partial m} = \alpha\delta > 0, \quad \frac{\partial a^*}{\partial m} = -\frac{\phi}{\gamma} < 0.$$

2.2. Etapa 1: problema del Profesor

Reemplazando $e^*(m) = \alpha(1 + \delta m)$ en (2),

$$U^P(m) = \alpha(1 + \delta m) - \frac{\kappa}{2} m^2.$$

Entonces $\frac{d^2 U^P}{dm^2} = -\kappa < 0$, por lo que U^P es estrictamente cóncava y la CPO es necesaria y suficiente:

$$\frac{dU^P}{dm} = \alpha\delta - \kappa m = 0 \implies m^* = \frac{\alpha\delta}{\kappa}. \quad (4)$$

Como $\alpha, \delta, \kappa > 0$, $m^* > 0$; y $m^* \leq 1 \iff \alpha\delta \leq \kappa$. Si $\alpha\delta > \kappa$ el óptimo es la esquina $m^* = 1$.

Bajo los Supuestos 1 y las condiciones de interioridad $m^* \leq 1$ y $m^* \leq \alpha/\phi$, el juego Γ tiene un único equilibrio perfecto en subjuegos

$$(m^*, \sigma^*(\cdot)), \quad \sigma^*(m) = (e^*(m), a^*(m)),$$

con $e^*(m), a^*(m)$ dados por (3) y m^* por (4). El equilibrio es

$$m^* = \frac{\alpha\delta}{\kappa}, \quad e^*(m^*) = \alpha\left(1 + \frac{\alpha\delta^2}{\kappa}\right), \quad a^*(m^*) = \frac{1}{\gamma}\left(\alpha - \frac{\alpha\delta\phi}{\kappa}\right).$$

Aquí la estática comparativa del equilibrio es la siguiente:

$$\frac{\partial m^*}{\partial \alpha} = \frac{\delta}{\kappa} > 0, \quad \frac{\partial m^*}{\partial \delta} = \frac{\alpha}{\kappa} > 0, \quad \frac{\partial m^*}{\partial \kappa} = -\frac{\alpha\delta}{\kappa^2} < 0, \quad \frac{\partial m^*}{\partial \phi} = \frac{\partial m^*}{\partial \gamma} = 0.$$

Esto es, en el modelo base la severidad ϕ no afecta el monitoreo óptimo, aparece solo en a^* y el Profesor valora únicamente $f(e) = e$, que es independiente de ϕ .

3. Extensión: castigo sobre el gap

La extensión reemplaza el castigo proporcional a la delegación, $m\phi a$, por un castigo proporcional al diferencial entre delegación y esfuerzo, valorado en unidades de nota:

$$U^S(e, a | m) = \alpha(e + a) - \alpha\phi m(a - e) - \frac{e^2}{2(1 + \delta m)} - \frac{\gamma}{2}a^2, \quad \phi \in [0, 1]. \quad (5)$$

El pago del Profesor mantiene la forma (2).

3.1. Etapa 2: problema del Estudiante

La función de utilidad sigue siendo estrictamente cóncava. Simplemente se añadió el término $-\alpha\phi m(a - e)$, al ser lineal en (e, a) no se alteran las segundas derivadas y el Hessiano es el mismo que el caso base. Las CPO son :

$$\frac{\partial U^S}{\partial e} = \alpha + \alpha\phi m - \frac{e}{1 + \delta m} = 0, \quad \frac{\partial U^S}{\partial a} = \alpha - \alpha\phi m - \gamma a = 0.$$

De la primera, $e^*(m) = \alpha(1 + \phi m)(1 + \delta m)$; de la segunda, $a^*(m) = \alpha(1 - \phi m)/\gamma$.

La mejor respuesta interior del Estudiante es

$$e^*(m) = \alpha(1 + \phi m)(1 + \delta m), \quad a^*(m) = \frac{\alpha(1 - \phi m)}{\gamma}. \quad (6)$$

Se siguen manteniendo los supuestos aplicados a las relaciones entre parámetros. A diferencia del modelo base, ϕ entra ahora en e^* , $\partial e^*/\partial \phi = \alpha m(1 + \delta m) \geq 0$. El término lineal aporta $+\alpha\phi m$ al retorno marginal del esfuerzo.

3.2. Etapa 1: problema del Profesor

Vamos a asumir que

$$\kappa > 2\alpha\phi\delta \quad (7)$$

Sustituyendo $e^*(m)$ de (6) en (2),

$$U^P(m) = \alpha(1 + \phi m)(1 + \delta m) - \frac{\kappa}{2}m^2 = \alpha + \alpha(\phi + \delta)m + \left(\alpha\phi\delta - \frac{\kappa}{2}\right)m^2.$$

La segunda derivada es

$$\frac{d^2U^P}{dm^2} = 2\alpha\phi\delta - \kappa,$$

que es negativa si y solo si $\kappa > 2\alpha\phi\delta$. Esto es, bajo (7) U^P es estrictamente cóncava. La CPO es:

$$\frac{dU^P}{dm} = \alpha(\phi + \delta) + (2\alpha\phi\delta - \kappa)m = 0 \implies m^* = \frac{\alpha(\phi + \delta)}{\kappa - 2\alpha\phi\delta}.$$

Como el numerador es positivo y, por (7), el denominador también, $m^* > 0$.

En conclusión, bajo (7), U^P es estrictamente cóncava en m y se maximiza en

$$m^* = \frac{\alpha(\phi + \delta)}{\kappa - 2\alpha\phi\delta}. \quad (8)$$

3.3. Supuestos de regularidad de la extensión

Para la extensión se requieren más supuestos, no se añade mucho valor a las recomendaciones de política por lo que el modelo base es nuestro modelo canónico.

Para que (8) sea el EPS interior y que sea interpretable como castigo se requiere:

Supuesto 2 (Concavidad). $\kappa > 2\alpha\phi\delta$, lo cual garantiza concavidad y por lo tanto máximo.

Supuesto 3 (Monitoreo interior). $m^* \leq 1$, dado que es una probabilidad, equivalentemente $\kappa \geq \alpha(\phi + \delta) + 2\alpha\phi\delta = \alpha(\phi + \delta + 2\phi\delta)$.

Supuesto 4 (Delegación no negativa). $a^*(m^*) \geq 0 \iff \phi m^* \leq 1$. Lo cual se satisface si $\phi \in [0, 1]$ y $m^* \leq 1$.

Supuesto 5 (Gap no negativo). $a^*(m^*) \geq e^*(m^*)$, supuesto fuerte.

El supuesto 5 es necesario para que siempre sea un posible castigo y no un subsidio a esforzarse más que delegar. Si se cumple cierta condición entonces el supuesto se cumple como consecuencia. Esto quiere decir que si

$$\frac{1 - \phi}{\gamma} \geq (1 + \phi)(1 + \delta), \quad (9)$$

entonces $a^*(m) \geq e^*(m)$ para todo $m \in [0, 1]$, y en particular en m^* .

Se puede mostrar de la siguiente manera, dado (6), $a^*(m) = \frac{\alpha(1-\phi m)}{\gamma}$ es decreciente en m y $e^*(m) = \alpha(1 + \phi m)(1 + \delta m)$ es creciente en m . Luego la diferencia $a^*(m) - e^*(m)$ es decreciente en m y alcanza su mínimo sobre $[0, 1]$ en $m = 1$. Evaluando en $m = 1$:

$$a^*(1) \geq e^*(1) \iff \frac{\alpha(1 - \phi)}{\gamma} \geq \alpha(1 + \phi)(1 + \delta) \iff \frac{1 - \phi}{\gamma} \geq (1 + \phi)(1 + \delta),$$

que es (9). Si esta se cumple, $a^*(m) - e^*(m) \geq a^*(1) - e^*(1) \geq 0$ en todo $[0, 1]$.

Dados los anteriores supuestos y condiciones la extensión tiene un EPS único dado por $(m^*, \sigma^*(\cdot))$ con $\sigma^*(m) = (e^*(m), a^*(m))$.

3.4. Estática comparativa de la extensión

En el equilibrio interior (8), bajo el Supuesto 2,

$$\frac{\partial m^*}{\partial \phi} = \frac{\alpha(\kappa + 2\alpha\delta^2)}{(\kappa - 2\alpha\phi\delta)^2} > 0, \quad \frac{\partial m^*}{\partial \delta} = \frac{\alpha(\kappa + 2\alpha\phi^2)}{(\kappa - 2\alpha\phi\delta)^2} > 0.$$

En particular, a diferencia del modelo base, $\partial m^*/\partial \phi > 0$. La derivada parcial es compleja, pero se puede mostrar de la siguiente manera: Definiendo $m^* = N/D$ con $N = \alpha(\phi + \delta)$ y $D = \kappa - 2\alpha\phi\delta$. Para ϕ : $\partial_\phi N = \alpha$ y $\partial_\phi D = -2\alpha\delta$, de modo que

$$\frac{\partial m^*}{\partial \phi} = \frac{(\partial_\phi N)D - N(\partial_\phi D)}{D^2} = \frac{\alpha D + 2\alpha\delta \alpha(\phi + \delta)}{D^2} = \frac{\alpha(\kappa - 2\alpha\phi\delta) + 2\alpha^2\delta(\phi + \delta)}{D^2}.$$

El numerador es $\alpha\kappa - 2\alpha^2\phi\delta + 2\alpha^2\phi\delta + 2\alpha^2\delta^2 = \alpha(\kappa + 2\alpha\delta^2)$, de donde la primera igualdad, es positivo bajo el Supuesto 1. El cálculo para δ es simétrico ($\partial_\delta N = \alpha$, $\partial_\delta D = -2\alpha\phi$) y da $\alpha(\kappa + 2\alpha\phi^2)$ en el numerador.