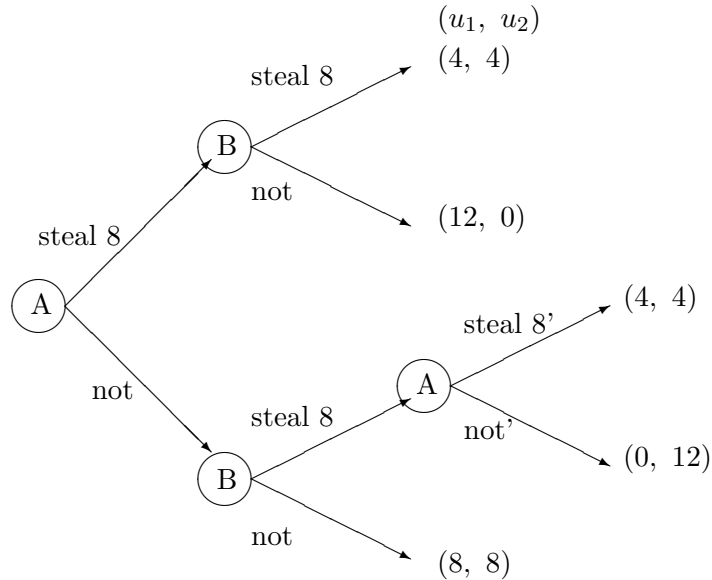


2025 年度 ゲーム理論 a 演習第 2 回解答

Takako Fujiwara-Greve

1. (人生はお話で与えられるので、お話からゲームを作れるようになります。)

(a) 樹形図は以下。



(b) A の情報集合は始点と歴史 (not, steal 8) 後にあるので、この順番に行動計画を書くと、純戦略は

(steal 8, steal 8'), (steal 8, not'), (not, steal 8'), (not, not')

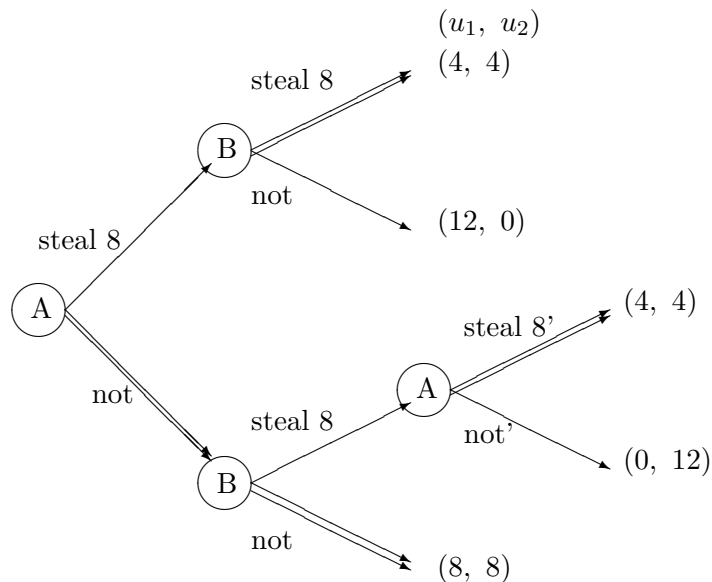
の 4 つである。

(c) B の情報集合は A が Step 1 で steal 8 を選んだ後と Step 1 で not を選んだ後の 2 つであり、この順番に行動計画を書くと、純戦略は

(steal 8, steal 8), (steal 8, not), (not, steal 8), (not, not)

の 4 つである。(樹形図で、B の異なる情報集合における行動の名前を変えていたら、それに合わせて変えること。ここでは変えていないので同じ名前になっているだけ。)

(d) ただ一つある。完全情報なので、後ろ向きの帰納法でも部分ゲーム完全均衡でも同じ解き方となる。最後の情報集合から、最適な行動を 2 重の矢印で描くと以下ようになる。



したがって、純戦略による解（部分ゲーム完全均衡）はただ一つで、それは

((not, steal 8'), (steal 8, not))

という組み合わせである。（A の戦略は行動の名前にダッシュがついている部分が正確に書けなければならない。）

2. (a) 純戦略のナッシュ均衡は2つあって、(B,C) と (B,R) である。

また、(i)P2 にとって L は R に厳密に支配されているので、ナッシュ均衡では行われぬ。

(ii)P1 が B という純戦略をしているとき以外は P2 の最適反応は R のみである。(iii)P1 が B という純戦略をしているときは、P2 の最適反応全体の集合は

$$\{qC + (1 - q)R \mid 0 \leq q \leq 1\}$$

である。

そこで、P2 の（ナッシュ均衡で行いうる）混合戦略を $\sigma_2 = qC + (1 - q)R$ とすると、P1 の最適反応は q にのみ依存しているとしてよく、

$$BR_1(q) = B, \forall q \in [0, 1]$$

である。

以上から、混合戦略のナッシュ均衡全体の集合は

$$\{(B, qC + (1 - q)R) \mid 0 \leq q \leq 1\}$$

である。

- (b) ない。純戦略に限っているので、2 回目は (B,C) または (B,R) でなければならないが、いずれにせよ P2 の利得は 3 である。したがって 1 回目に何をしようとも P2 の 2 回目の利得を変えることはできないから、1 回目に P2 が最適反応をしている行動の組み合わせしか部分ゲーム完全均衡にならない。

- (c) グリム・トリガー戦略を作ることを考える。 δ の範囲を大きくするためには、2 人への罰ができる限り厳しくなるようにする。しかし、このゲームでは、P1 は (T,L) から逸脱する理由はないので実質的には P2 が逸脱したときに最も厳しい罰を純戦略の範囲で考えればよい。それは (B,C) または (B,R) のどちらでもよいので、以下では例として (B,C) を罰としたグリム・トリガー戦略を書いておく。

P1 の $G^\infty(\delta)$ における純戦略 $s_1 = (x_1, x_2, \dots)$ は、第 1 期の行動は $x_1 = T$ 、その後の任意の $t \geq 2$ については

$$x_t(h_t) = \begin{cases} T & \text{if } h_t = (T, L)^{t-1} \\ B & \text{otherwise} \end{cases}$$

という関数 $x_t: \{T, B\} \times \{L, C, R\}^{t-1} \rightarrow \{T, B\}$ になる行動計画 $(x_1, x_2, \dots)$ とする。

P2 の純戦略 $s_2 = (y_1, y_2, \dots)$ は、第 1 期の行動は $y_1 = L$ 、その後の任意の $t \geq 2$ については

$$y_t(h_t) = \begin{cases} L & \text{if } h_t = (T, L)^{t-1} \\ C & \text{otherwise} \end{cases}$$

という関数になる行動計画 $(y_1, y_2, \dots)$ とする。

（注：長期戦略の書き方は読者にわかればよい。

例えば $\{T, B\} \times \{L, C, R\}^0 = \emptyset$ と定義しておくとも第 1 期における空の歴史も含めて任意の t 期における歴史を $h_t \in \{T, B\} \times \{L, C, R\}^{t-1}$ と書けるので、第 1 期と

それ以降を分けなくてもよくて、戦略の定義域を $\cup_{t=1,2,\dots}[\{T, B\} \times \{L, C, R\}]^{t-1}$ に拡張して、

$$s_1(h_t) = \begin{cases} T & \text{if } h_t = (T, L)^{t-1} \\ B & \text{otherwise} \end{cases}, \quad s_2(h_t) = \begin{cases} L & \text{if } h_t = (T, L)^{t-1} \\ C & \text{otherwise} \end{cases}$$

という関数である、と書いてもよい。ただし、読者にわかるように書くこと。）

上記の戦略の組み合わせ (s_1, s_2) が部分ゲーム完全均衡になるような δ の範囲を求める。まず、 (T, L) が続いている任意の歴史の後の部分ゲームにおいては (s_1, s_2) は段階ゲームのナッシュ均衡である (B, C) をその後ずっと行うことになるので、任意の δ についてそのような継続戦略の組み合わせはナッシュ均衡である。

そこで、第 1 期またはこれまではずっと (T, L) だった歴史の後の部分ゲームのみ考えればよい。P1 の観点からは、P2 の継続戦略は

$$y_t(h_t) = \begin{cases} L & \text{if } h_t = (T, L)^{t-1} \\ C & \text{otherwise} \end{cases}$$

という形の関数が続いているものである。ここで P1 が s_1 からの one-step deviation を行なってもそもそも段階ゲームで 5 より高い利得を得られないので、割引総利得が高まることはない。したがってこれらの部分ゲームにおいて s_1 のその部分ゲームに制限した継続戦略に従うことは最適反応である。

P2 の観点からは s_2 の継続戦略から one-step deviation を行なって R をすると 1 期間だけ 6 を得られるが、その後の部分ゲームにおいて P1 の継続戦略はずっと B を行うというものになるので、one-step deviation から得られる最大の割引総利得は

$$6 + \delta \cdot 3 + \delta^2 \cdot 3 + \dots = 6 + \frac{3\delta}{1-\delta}$$

である。これに対し、 s_2 の継続戦略にずっと従っていれば

$$5 + \delta \cdot 5 + \delta^2 \cdot 5 + \dots = \frac{5}{1-\delta}$$

が得られるので、

$$\frac{5}{1-\delta} \geq 6 + \frac{3\delta}{1-\delta} \iff \delta \geq \frac{1}{3}$$

ならば s_2 の継続戦略に従うことが最適反応である。

以上から $[\frac{1}{3}, 1)$ が $G^\infty(\delta)$ において每期 (T, L) を行う純戦略の部分ゲーム完全均衡が存在する最大の範囲である。

3. X 社は図にもあるように情報集合が二つあり、(新製品があるとき、ないとき)の順に書くと、純戦略の集合は $S_X = \{(\text{Sale}, \text{Sale}'), (\text{Sale}, \text{Not}'), (\text{Not}, \text{Sale}'), (\text{Not}, \text{Not}')\}$ である。

情報のない Y 社は実質的にはただの同時ゲームをしているので純戦略の集合は $S_Y = \{20, 10, \text{Not}\}$ である。