



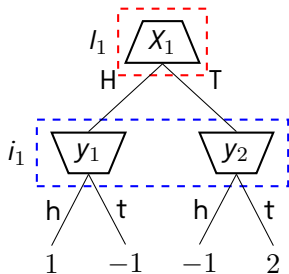
Solving Imperfect Information EFGs

Ondřej Kubíček

Artificial Intelligence Center
Faculty of Electrical Engineering
Czech Technical University in Prague

October 22, 2024

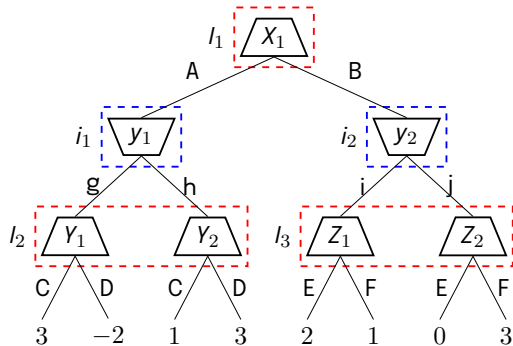
Problems with solving imperfect information games



- First player in I_1 chooses between H and T based on the strategy that the player 2 plays.
- Second player in i_1 chooses between h and t based on the probability if it is in y_1 or y_2 , that directly corresponds to the strategy of player 1.
- Backward induction will not work because of this interconnected dependency.
- Generally even in perfect recall imperfect information games, the policy depends on both policy in previous and subsequent parts of the game tree

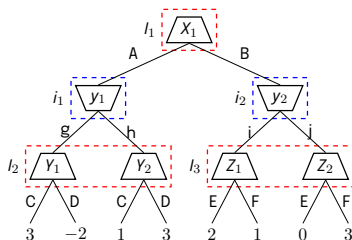
Example

- Induced normal-form game can include the same leaf multiple times.
- Playing some actions may invalidate different actions in future.



	gi	gj	hi	hj
ACE	3	3	1	1
ACF	3	3	1	1
ADE	-2	-2	3	3
ADF	-2	-2	3	3
BCE	2	0	2	0
BCF	1	3	1	3
BDE	2	0	2	0
BDF	1	3	1	3

- Ordered list of all the actions that may be played in a single playthrough of the game is called Sequence.
- We denote all possible sequences of player i as Σ_i .



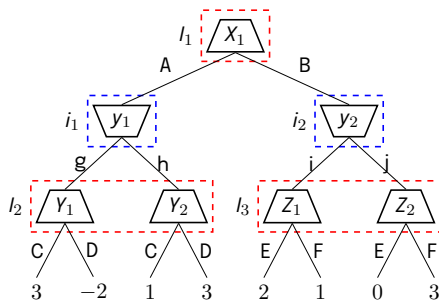
Σ_1	Σ_2
$\emptyset$	$\emptyset$
A	g
B	h
AC	i
AD	j
BE	
BF	

- Realization plan $r_i(\sigma_i)$ is a probability that σ_i will be played, assuming that the other player plays only actions that allow σ_i to be executed
- Let us assume that σ_i leads to the info set I_i . Behavioral strategy $\pi(I_i, a)$ of playing action a in this info set is computed as

$$\pi(I_i, a) = \begin{cases} \frac{r(\sigma_i a)}{r(\sigma_i)} & \text{if } r(\sigma_i) > 0 \\ 0 & \text{otherwise} \end{cases}$$

- $\sigma_i a$ represents a extensions of sequence σ_i with action a .

Example

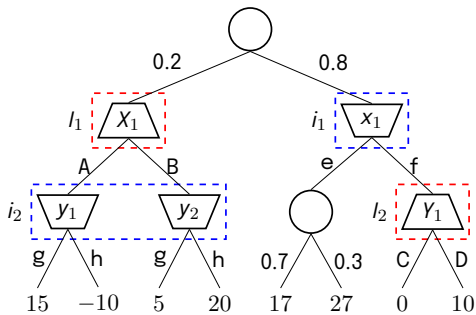


Σ_1	Σ_2
$\emptyset$	$\emptyset$
A	g
B	h
AC	i
AD	j
BE	
BF	

- $r_1(\emptyset) = 1$
- $r_1(A) + r_1(B) = r_1(\emptyset)$
- $r_1(AC) + r_1(AD) = r_1(A)$
- $r_1(BE) + r_1(BF) = r_1(B)$

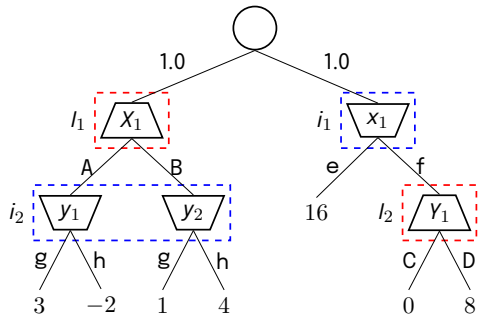
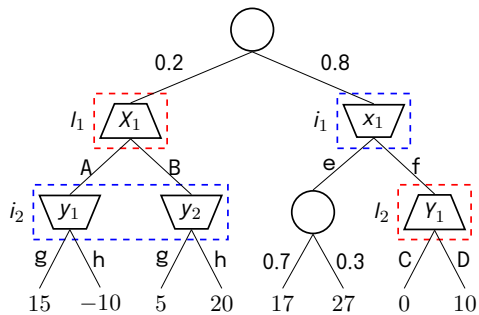
- $r_2(\emptyset) = 1$
- $r_2(g) + r_2(h) = r_2(\emptyset)$
- $r_2(i) + r_2(j) = r_2(\emptyset)$

Propagating chance node



- Chance nodes have fixed probabilities in the game
- These probabilities can be propagated from the chance node to the terminal utilities
- Chance nodes that do not reveal any information until the end of the game can be pruned away completely

Example



- Extended utility function for sequences is $g : \Sigma_1 \times \Sigma_2 \rightarrow \mathbb{R}$

$$g(\sigma_1, \sigma_2) = \sum_{z \in \mathcal{Z}'} \mathcal{C}(z) u(z)$$

- $\mathcal{C}(z)$ is a probability that leaf z was reached due to chance nodes along the way.
- $\mathcal{Z}' \subseteq \mathcal{Z}$ are all the terminal histories that could be reached by sequences σ_1, σ_2

Example

$$g(\emptyset, \emptyset) = 0$$

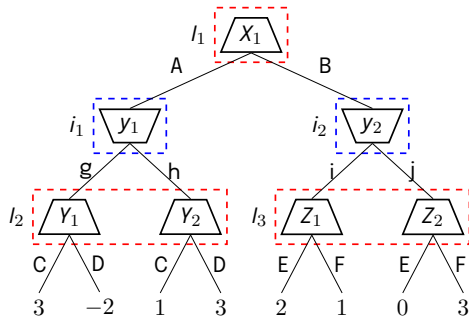
$$g(AC, i) = 0$$

$$g(BF, j) = 3$$

$$g(AD, \emptyset) = 0$$

$$g(\emptyset, g) = 0$$

$$g(A, g) = 0$$



Σ_1	Σ_2
$\emptyset$	$\emptyset$
A	g
B	h
AC	i
AD	j
BE	
BF	

Sequence-form Linear Program (SQF)

$$\max_{r_1, v} v(l_{\text{root}}) \quad (1)$$

$$\text{s.t. } r_1(\emptyset) = 1 \quad (2)$$

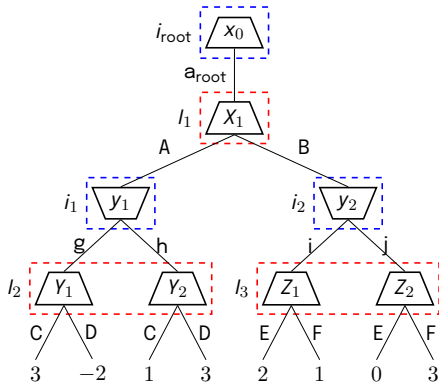
$$r_1(\sigma_1) \geq 0 \quad \forall \sigma_1 \in \Sigma_1 \quad (3)$$

$$\sum_{a \in \mathcal{A}_1(l_1)} r(\sigma_1 a) = r_1(\sigma_1) \quad \forall l_1 \in \mathcal{I}_1, \sigma_1 = s_1(l_1) \quad (4)$$

$$\sum_{l'_2 \in \mathcal{I}_2: \sigma_2 a = s_2(l'_2)} v(l'_2) + \sum_{\sigma_1 \in \Sigma_1} g(\sigma_1, \sigma_2 a) r_1(\sigma_1) \geq v(l_2) \quad \forall l_2 \in \mathcal{I}_2, \sigma_2 = s_2(l_2), \forall a \in \mathcal{A}(l_2) \quad (5)$$

- Variables are realization plans r_1 for all the sequences of player 1 and expected values v for each opponents info set, if it plays a best response.
- s_i returns for a given info set l_i a sequence σ_i that leads to this info set.

Example



$$\max_{r_1, v}(v(i_{\text{root}}))$$

$$r_1(\emptyset) = 1$$

$$r_1(A) + r_1(B) = r_1(\emptyset)$$

$$r_1(AC) + r_1(AD) = r_1(A)$$

$$r_1(BE) + r_1(BF) = r_1(B)$$

$$v(i_1) + v(i_2) \geq v(i_{\text{root}})$$

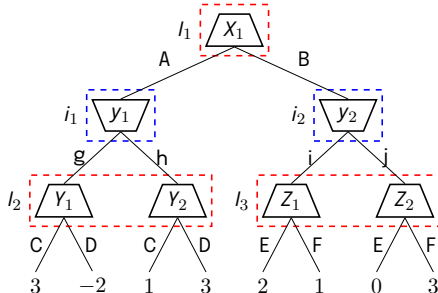
$$3r_1(AC) - 2r_1(AD) \geq v(i_1)$$

$$1r_1(AC) + 3r_1(AD) \geq v(i_1)$$

$$2r_1(BE) + 1r_1(BF) \geq v(i_2)$$

$$0r_1(BE) + 3r_1(BF) \geq v(i_2)$$

Example



$$\min_{r_2, v}(v(l_1))$$

$$r_2(\emptyset) = 1$$

$$r_2(g) + r_2(h) = r_2(\emptyset)$$

$$r_2(i) + r_2(j) = r_2(\emptyset)$$

$$v(l_2) \leq v(l_1)$$

$$v(l_3) \leq v(l_1)$$

$$3r_2(g) + 1r_2(h) \leq v(l_2)$$

$$-2r_2(g) + 3r_2(h) \leq v(l_2)$$

$$2r_2(i) + 0r_2(j) \leq v(l_3)$$

$$1r_2(i) + 3r_2(j) \leq v(l_3)$$

- The fastest exact algorithm
- Easy to implement
- Poor scaling due to memory requirements of the linear program
- Hard to fine-tune for specific domains to increase performance
- Cannot be used to solve the game only partially

- Not all strategies in the game are necessary to find Nash equilibrium.
- Double oracle creates smaller game with less strategies and iteratively adds new strategies until it finds Nash equilibrium of the underlying game.
- In the worst case scenario the double oracle algorithm has to add all the strategies from the original game.
- However, in most cases, the restricted game is much smaller than the original game.

- Using sequences instead of pure strategies is more complicated.
- With limited amount of sequences, some reachable parts of the game tree may not have any sequence to be played.
- To avoid this, in each information set, there is some default action that should be played.
- Default action does not have to be defined explicitly.
- Instead of keeping the full tree, the nodes with default action can be replaced by terminal node.
- The value of this terminal node corresponds to the value if opponent picks best response against the default action.

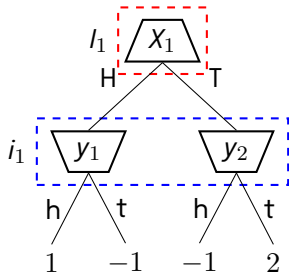
- Can solve larger games than SQF.
- Without any additional information the algorithm identifies, which strategies are not important for good solution.
- Computing best response is faster than solving the Linear program and can be improved with some heuristics for specific problems.
- Harder to implement, due to the need to construct the valid restricted game.
- Still requires SQF to solve the restricted game, which is the primary limitation of the method.
- In games where all sequences have to be considered, it is slower than the SQF.

- Framework that generalizes double oracle algorithm with empirical game theoretical analysis.
- It was designed to use double oracle with deep learning with application to very large games.
- Instead of using pure strategies, it uses the full policies.
- The utility of 2 policies against each other can be computed by just letting them play in heads-to-heads.
- This creates matrix game that can be solved through usual means
- Best response can be computed through reinforcement learning, since policy of 1 player is fixed.

- Mixed strategies are not well suited for solving the imperfect information EFGs.
- Sequences can be used only in perfect recall games, but avoid necessity to define each action for each information set.
- Probability from chance nodes can be propagated up to the leaf utilities.
- If chance node does not reveal any information until the end of the game, it can be removed from the game completely.
- Sequence-form linear program for solving EFGs expands the ideas from linear program for normal-form games to the extensive-form setting.
- It uses the realization plans instead of mixed strategies and expected values for each opponent's info set.
- Double oracle algorithm can be generalized to a very large games.

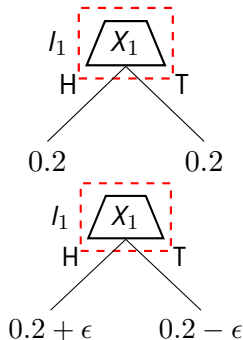
- What if we do not want to solve the whole game, but we only want the strategy in some part of the game.
- In perfect information games, this is done by Monte-Carlo Tree Search.
- What about imperfect information games?
- There are several problems not present in perfect information games.
- One is that the algorithm cannot assume the current state of the game since it is not observable.
- Least amount of considered states has to be closed on the information known by all players.

Appendix: Depth-limited solving



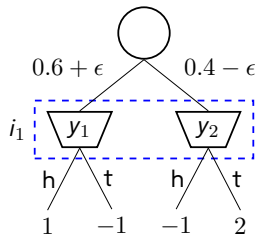
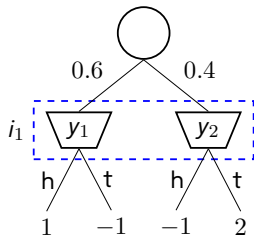
- Optimal strategy in this game is by both players playing heads with 60%.
- What if we are first player and we want to solve this game with depth-limit?
- What if we are the second player and we want to solve this game from our decision onward?
- Let us assume we know opponent's optimal strategy.

Appendix: Depth-limited solving



- We assume optimal opponent.
- Value of playing heads is $0.6 \cdot 1 - 0.4 \cdot 1 = 0.2$
- Value of playing tails is $-0.6 \cdot 1 + 0.4 \cdot 2 = 0.2$
- Is solution of such a game optimal in the original game?
- What happens with the solution if $\epsilon > 0$ or $\epsilon < 0$?

Appendix: Depth-limited solving



- We still assume optimal opponent
- Value of playing heads is $-0.6 \cdot 1 + 0.4 \cdot 1 = -0.2$
- Value of playing tails is $0.4 \cdot 1 - 0.6 \cdot 2 = -0.2$
- What would be a solution of such a game?
- What happens with the solution if $\epsilon > 0$ or $\epsilon < 0$?

- Both of these problems happen because of the indifference principal.
- In order to have safe depth-limited solving these has to be addressed.
- One possible solution to the first problem is to allow players use several strategies in the future.
- Second is to use iterative approach to converge to a Nash, where after the depth-limit the value depends on the policy in preceding parts
- Second problem is solved by allowing opponent to pick any strategy in the previous part of the game.