

Approachability, Calibration, and Adaptive Algorithms

Duarte Gonçalves

University College London

Topics in Economic Theory

Overview

1. Learning in Games
2. Approachability
3. Calibration
4. Adaptive Algorithms

Overview

1. Learning in Games

2. Approachability

3. Calibration

4. Adaptive Algorithms

How do people get to play equilibrium?

Main question of interest in 'learning in games' ($\neq$ games with learning)

Goals

Provide foundations for existing equilibrium concepts.

Capture lab behaviour.

Predict adjustment dynamics transitioning to new equilibrium.

(akin to 'impulse response' in macro; uncommon but definitely worth investigating)

Select equilibria.

Algorithm to solve for equilibria.

Explain persistence of heuristics/non-equilibrium behaviour.

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Approachability with Time-Varying Payoffs
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Approachability

Setup:

Two players face repeated game. m -dimensional goal/multi-utility representation.

Goal of each player: control average vector of attributes s.t. approaches/excludes (keeps away) target set $S \subset \mathbb{R}^m$.

Brief detour: rationalising multi-utility.

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Approachability with Time-Varying Payoffs
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Multi-Utility Representation

$\succsim \subseteq X^2$: preorder (reflexive, transitive).

Want to allow for incompleteness.

x and y are **incomparable** if $\neg(x \succsim y \text{ or } y \succsim x)$. Ow, they are comparable.

E.g., choice by unanimity, incomparable attributes.

Definition

For any binary relation $\succsim$ on X with symmetric part $\sim$, for any $x \in X$, x 's **equivalence class** is $[x] := \{y \in X \mid x \sim y\}$ and the set of equivalence classes $\hat{X} := \{[x], x \in X\}$.

Remark

For any preorder $\succsim$ on X , let $\hat{\succsim}$ on $\hat{X} : \forall x, y \in X : [x] \hat{\succsim} [y] \text{ if } x \succsim y$. Then, $\hat{\succsim}$ is partial order.

Multi-Utility Representation

$\succsim \subseteq X^2$: preorder (reflexive, transitive).

Want to allow for incompleteness.

x and y are **incomparable** if $\neg(x \succsim y \text{ or } y \succsim x)$. Otherwise, they are comparable.

E.g., choice by unanimity, incomparable attributes.

Set of all linear extensions of $\hat{\succsim}$ on X : $\mathcal{L}(\hat{X}, \hat{\succsim})$.

Szpilrajn's Theorem: any partial order can be extended to a linear order.

Remark 1: $\implies \mathcal{L}(\hat{X}, \hat{\succsim}) \neq \emptyset$.

Remark 2: $\hat{\succsim} = \bigcap_{\succsim \in \mathcal{L}(\hat{X}, \hat{\succsim})} \succsim$.

Order dimension: $\dim(X, \succsim) := \min\{k \in \mathbb{N} \mid \exists \succsim_i \in \mathcal{L}(\hat{X}, \hat{\succsim}), i = 1, \dots, k : \hat{\succsim} = \bigcap_{i=1}^k \succsim_i\}$.

$\dim(X, \succsim)$: min number of linear extensions of $\hat{\succsim}$ whose intersection yields $\hat{\succsim}$.

Examples:

$\hat{\succsim}$ is linear order on X iff $\dim(X, \succsim) = 1$.

If no distinct x, y are comparable ($\hat{\succsim}$ is antichain) and $\dim(X, \succsim) = 2$ since

$$\hat{\succsim} = \geq \cap \leq.$$

If $X = 2^A$ and $|A| = \infty$, then $\dim(X, \subseteq) = \infty$.

Multi-Utility Representation

Definition

$\succsim \subseteq X^2$ admits a multi-utility representation $u : X \rightarrow \mathbb{R}^m$ with $m \in \mathbb{N}$ iff $\forall x, y \in X, x \succsim y \iff u(x) \geq u(y)$.

Proposition 1 (Ok 2002 JET)

Let $\succsim$ be preorder on X .

- (1) $\succsim$ admits a multi-utility representation u only if $\dim(X, \succsim) < \infty$.
- (2) If $\hat{X}$ countable, $\succsim$ admits a multi-utility representation u if and only if $\dim(X, \succsim) < \infty$.

Alternative (social) interpretation: $\exists U \subset \mathbb{R}^X$ such that $x \succsim y \iff u(x) \geq u(y) \forall u \in U$.

Proof Idea

Take X finite. Let $u_x(y) = \mathbf{1}_{\{y \succsim x\}}$. $u(y) = (u_x(y))_{x \in X}$.

Let $x \succsim y$. (a) $\forall z \in X : (u_z(x) = 0) \iff (z \not\succsim x) \implies (z \not\succsim y) \iff (u_z(y) = 0)$

(b) $\forall z \in X : (u_z(y) = 1) \iff (y \succsim z) \implies (x \succsim z) \iff (u_z(x) = 1)$.

(a) + (b) $\implies u(x) \geq u(y)$.

Multi-Utility Representation

Definition

$\succsim \subseteq X^2$ admits a multi-utility representation $u : X \rightarrow \mathbb{R}^m$ with $m \in \mathbb{N}$ iff $\forall x, y \in X, x \succsim y \iff u(x) \geq u(y)$.

Proposition 1 (Ok 2002 JET)

Let $\succsim$ be preorder on X .

- (1) $\succsim$ admits multi-utility representation u only if $\dim(X, \succsim) < \infty$.
- (2) If $\hat{X}$ countable, $\succsim$ admits a multi-utility representation u if and only if $\dim(X, \succsim) < \infty$.

Proposition 2 (Ok 2002 JET)

- (a) $X_0 = \times_{i=1}^k X_i$, with X_i be metric space and $\succsim_i$ be preorders on $X_i, i = 0, 1, \dots, k$;
- (b) Each X_i is s.t. $\{y_i \mid y_i \succ_i x_i\}$ is open for every $x_i \in X_i$ and $i = 1, \dots, k$; and
- (c) $x \succsim_0 y \iff x_i \succsim_i y_i \forall i = 1, \dots, k$.

If X_0 admits a countable $\succsim_0$ -dense subset, then $\succsim_0$ admits a multi-utility representation u which is continuous in the product topology.

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- **Approachability**
- Approachability with Time-Varying Payoffs
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Approachability

Setup:

Two players face repeated game. m -dimensional goal/multi-utility representation.

Goal of each player: control average vector of attributes s.t. approaches/excludes (keeps away) target set $S \subset \mathbb{R}^m$.

Back to approachability... Blackwell (1956). Good reference: Maschler, Solan, Zamir (2013 Book, Ch. 14).

Actions: A_i ; **Stage-Game Payoffs:** $u_1 : A_1 \times A_2 \rightarrow \mathbb{R}^m$ $u_2 := -u_1$; (endowed with $d(x, y) = \|x - y\|_2$);

Histories: $H_t := A^t$, $H := \cup_t H_t$; **Strategies:** $\sigma_i : H \rightarrow \Delta(A_i)$; $\lambda_i \in \Delta(A_i)$.

Expected Payoffs: $u_i(\lambda_i, \lambda_{-i}) = \sum_{a_i} \sum_{a_{-i}} \lambda_i(a_i) u_i(a_i, a_{-i}) \lambda_{-i}(a_{-i}) \in \mathbb{R}^m$.

Feasible Expected Payoffs for λ_i : $U_i(\lambda_i) := \{u_i(\lambda_i, \lambda_{-i}), \lambda_{-i} \in \Delta(A_{-i})\} \subseteq \mathbb{R}^m$.

Average Payoff: $\bar{u}_{i,t} = \frac{1}{t} \sum_{\ell=1}^t u_i(a_\ell)$.

Feasible Avg Payoffs: $\text{co}(u_i) := \text{co}(\{u_i(a), a \in A\}) \subseteq \mathbb{R}^m$.

Definition

$C \subseteq \mathbb{R}^m$ is

approachable by player i if $\exists \sigma_i$ s.t. $\forall \varepsilon > 0, \exists T : \forall \sigma_{-i}, \mathbb{P}^\sigma(d(\bar{u}_{i,T}, C) < \varepsilon, \forall t \geq T) > 1 - \varepsilon$; in this case, σ_i approaches C for player i ; and

excludable by player i if $\exists \delta$ s.t. set $C_\delta^C := \{x \mid d(x, C) \geq \delta\}$ is approachable by player i ; if strategy σ_i approaches C_δ^C , then it excludes C for player i .

Approachable by a player if can guarantee that average payoff approaches the set w.p.1 uniformly over opponent's strategies: $\mathbb{P}^\sigma(\lim_{t \rightarrow \infty} d(\bar{u}_t, C) = 0) = 1$.

Remark

- (1) If σ_i approaches (resp. excludes) C , then it approaches (resp. excludes) the closure of C .
- (2) C cannot be approachable by one player and excludable by the other.
- (3) If $C \subseteq D$ and C is approachable by a player, then D is approachable by the same player.

Definition

Let $D \subseteq \mathbb{R}^m$, $x, y \in \mathbb{R}^m$.

Normal cone of D at y is given by $N_D(y) := \{n \in \mathbb{R}^m \mid n \cdot (z - y) \leq 0, \forall z \in D\}$.

Projection of $x \in \mathbb{R}^m$ to D is given by $P_D(x) := \arg \min_{y \in D} \|y - x\|$.

Remark

(1) $(x - y) \in N_D(y) \implies y \in P_D(x)$.

(2) If D is convex, then $(x - y) \in N_D(y) \iff y \in P_D(x)$.

Definition

C is **B-set** for player i if, $\forall x \in \text{co}(u_i) \setminus C$, $\exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{U_i(\lambda_i)}(y)$, i.e., $\forall \lambda_{-i}$, $(u_i(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Definition

C is **B -set** for player i if, $\forall x \in \text{co}(u_i) \setminus C$, $\exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{U_i(\lambda_i)}(y)$,
i.e., $\forall \lambda_{-i}, (u_i(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Theorem

If C is closed B -set for player i for every t , then it is approachable by player i .

Convergence of Non-negative Almost Supermartingales

We'll need this:

Theorem (Robbins and Siegmund 1971)

Let $(\mathcal{F}_t)$ be filtration and $(V_t)_{t \geq 0}$ be nonnegative, adapted. Suppose there are nonnegative, $\mathcal{F}_t$ -adapted processes $(\xi_t), (\beta_t), (\zeta_t)$ s.t.

$$\mathbb{E}[V_{t+1} \mid \mathcal{F}_t] \leq (1 + \xi_t)V_t - \zeta_t + \beta_t, \quad t \geq 0,$$

with $\sum_{t=0}^{\infty} \xi_t < \infty$ and $\sum_{t=0}^{\infty} \beta_t < \infty$ a.s.

Then, V_t converges a.s. to a finite, nonnegative limit V_{∞} , and $\sum_{t=0}^{\infty} \zeta_t < \infty$ a.s.

Going beyond Doob's MCT: convergence for non-negative almost supermartingales.

Convergence of Non-negative Almost Supermartingales

Theorem (Robbins and Siegmund 1971)

Let $(\mathcal{F}_t)$ be filtration and $(V_t)_{t \geq 0}$ be nonnegative, adapted. Suppose there are nonnegative, $\mathcal{F}_t$ -adapted processes $(\xi_t), (\beta_t), (\zeta_t)$ s.t.

$$\mathbb{E}[V_{t+1} \mid \mathcal{F}_t] \leq (1 + \xi_t)V_t - \zeta_t + \beta_t, \quad t \geq 0,$$

with $\sum_{t=0}^{\infty} \xi_t < \infty$ and $\sum_{t=0}^{\infty} \beta_t < \infty$ a.s.

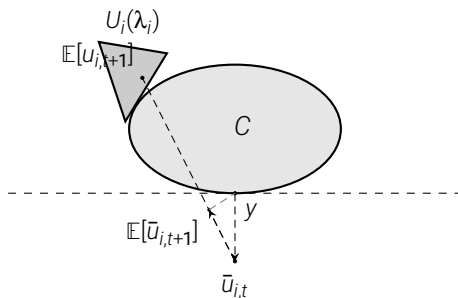
Then, V_t converges a.s. to a finite, nonnegative limit V_{∞} , and $\sum_{t=0}^{\infty} \zeta_t < \infty$ a.s.

Corollary

If nonnegative (V_t) satisfies $\mathbb{E}[V_t \mid \mathcal{H}_{t-1}] \leq (1 - \alpha_t)V_{t-1} + \beta_t$ with $\sum_t \alpha_t = \infty$ and $\sum_t \beta_t < \infty$, then $V_t \rightarrow 0$ a.s.

Useful corollary: $\xi_t = 0$, $\zeta_t = \alpha_t V_t$ with $\alpha_t \in [0, 1]$. If $\sum_t \alpha_t = \infty$, then $V_{\infty} = 0$ a.s.

Since $\sum_t \alpha_t V_t < \infty$; if $\sum_t \alpha_t = \infty$, only possible limit V_{∞} is 0.



Proof

Step 1: Projection Rule.

- (1) Let $y_{t-1} \in P_C(\bar{u}_{i,t-1})$.
- (2) If $\bar{u}_{i,t-1} \in C$, then play any $\lambda_{i,t} \in \Delta(A_i)$.
- (3) If $\bar{u}_{i,t-1} \notin C$, then, as C is B -set, there is $\lambda_{i,t} \in \Delta(A_i)$ s.t. $\forall \lambda_{-i}$,
 $(u_i(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) \leq 0$.
- (4) Play $\lambda_{i,t}$ at stage t .

Proof

Step 2: One-Step Inequality.

Update formula: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t}(u_{i,t}(a_t) - \bar{u}_{i,t-1})$. Let $V_t := d(\bar{u}_{i,t}, C) = \|\bar{u}_{i,t} - y_t\|^2$.

Since $y_{t-1} \in C$, then $\|\bar{u}_{i,t} - y_t\|^2 \leq \|\bar{u}_{i,t} - y_{t-1}\|^2$.

Expanding: $\|\bar{u}_{i,t} - y_{t-1}\|^2 = \|\bar{u}_{i,t-1} - y_{t-1} + \frac{1}{t}(u_{i,t}(a_t) - \bar{u}_{i,t-1})\|^2 = \|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \frac{2}{t}(u_{i,t}(a_t) - \bar{u}_{i,t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) + \frac{1}{t^2}\|u_{i,t}(a_t) - \bar{u}_{i,t-1}\|^2$.

- $(u_{i,t}(a_t) - \bar{u}_{i,t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) = (u_{i,t}(a_t) - y_{t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) - \|\bar{u}_{i,t-1} - y_{t-1}\|^2$.
- B-set: $\mathbb{E}[(u_{i,t}(a_t) - y_{t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) | H_{t-1}] = (u_{i,t}(\lambda_{i,t}, \lambda_{-i,t}) - \bar{u}_{i,t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) \leq 0$.
- $u_{i,t} \in [-M, M]^m \implies \|u_{i,t}(a_t) - \bar{u}_{i,t-1}\|^2 \leq 2\|u_{i,t}\|^2 \leq 4mM^2$.

$$\begin{aligned}
 \mathbb{E}[\|\bar{u}_{i,t} - y_t\|^2 | H_{t-1}] &\leq \mathbb{E}[\|\bar{u}_{i,t} - y_{t-1}\|^2 | H_{t-1}] \\
 &= \|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \mathbb{E}\left[\frac{1}{t}(u_{i,t}(a_t) - \bar{u}_{i,t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) + \frac{1}{t^2}\|u_{i,t}(a_t) - \bar{u}_{i,t-1}\|^2 \mid H_{t-1}\right] \\
 &\leq \|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \frac{2}{t}\mathbb{E}[(u_{i,t}(a_t) - y_{t-1}) \cdot (\bar{u}_{i,t-1} - y_{t-1}) \mid H_{t-1}] - \frac{1}{t}\|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \frac{1}{t^2}4mM^2 \\
 &\leq \left(1 - \frac{2}{t}\right) \|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \frac{1}{t^2}4mM^2
 \end{aligned}$$

Proof

Step 3: Apply Robbins and Siegmund (1971).

$$\mathbb{E}[\|\bar{u}_{i,t} - y_t\|^2 \mid H_{t-1}] \leq \left(1 - \frac{2}{t}\right) \|\bar{u}_{i,t-1} - y_{t-1}\|^2 + \frac{1}{t^2} 4mM^2$$

Satisfies conditions of Robbins and Siegmund (1971)'s corollary.

$$\implies \|\bar{u}_{i,t} - y_t\|^2 \rightarrow 0 \text{ a.s.}$$

Theorem (Blackwell 1956)

If C is B -set for player i , then it is approachable by player i .

Generalisations and Variations:

Lehrer (2002 IJGT): generalises Blackwell's approachability theorem to infinite-dimensional spaces.

Hou (1971 AMS): A closed set C is approachable by player i if and only if it contains a B -set for player i .

Theorem 14.25 (Maschler, Solan, and Zamir 2013)

If C is convex and closed, then C is approachable by one player if and only if it is not excludable by the other player.

Notation:

$H(x, y) := \{z \in \mathbb{R}^m \mid (x - y) \cdot (z - y) = 0\}$ Hyperplane passing through y that's orthogonal/perpendicular to line passing through x and y .

$H^-(x, y) := \{z \in \mathbb{R}^m \mid (x - y) \cdot (z - y) \leq 0\}$ Half-space defined by hyperplane $H(x, y)$.

Theorem 14.24 (Maschler, Solan, and Zamir 2013)

If C is convex and closed, then the following are equivalent:

- (1) C is approachable by player i ; (2) C is B -set for player i ; and
- (3) $\forall$ half-spaces $H^-(x, y)$ containing C , $\exists \lambda_i \in \Delta(A_i) : \forall \lambda_{-i}, u_i(\lambda_i, \lambda_{-i}) \in \subseteq H^-(x, y)$.

For convex sets, approachability is equivalent to B -set property!

Since any half-space containing C is approachable, closed and convex, it must also be a B -set. [(3) says a bit more than this, but this is the idea.]

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Approachability with Time-Varying Payoffs
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Approachability with Time-Varying Payoffs

Setup:

Actions A_i , $i = 1, 2$, $A := A_1 \times A_2$; Histories $H_t := A^t$, $H := \cup_t H_t$; Strategies $\sigma_i : H \rightarrow \Delta(A_i)$, $\lambda_i \in \Delta(A_i)$.

For every $t = 1, 2, \dots$, $i = 1, 2$, $u_{i,t} : A \rightarrow [-M, M]^m$ with $u_{2,t} = -u_{1,t}$.

$$U_{i,t}(\lambda_i) := \{u_{i,t}(\lambda_i, \lambda_{-i}), \lambda_{-i} \in \Delta(A_{-i})\}.$$

(Note uniform bound on payoffs.)

$$\text{Let } \bar{u}_{i,T} := \frac{1}{T} \sum_{t=1}^T u_{i,t}(a_t).$$

Definition

$C \subseteq \mathbb{R}^m$ is **approachable** by player i if $\exists \sigma_i$ s.t. $\forall \epsilon > 0$, $\exists T : \forall \sigma_{-i}$, $\mathbb{P}^\sigma(d(\bar{u}_{i,t}, C) < \epsilon, \forall t \geq T) > 1 - \epsilon$; in this case, σ_i approaches C for player i .

Approachability with Time-Varying Payoffs

Definition

C is strong **B-set** for player i if, $\forall t, \forall x \in \bigcup_t \text{co}(u_{i,t}) \setminus C, \exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{U_{i,t}(\lambda_i)}(y)$, i.e., $\forall \lambda_{-i}, (u_{i,t}(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Theorem

If C is closed strong B -set for player i , then, C is approachable by player i .

Proof we used can be used with minor modifications.

Approachability with Time-Varying Payoffs

Definition

C is **B -set** for player i at time t if, $\forall x \in \text{co}(u_{i,t}) \setminus C$, $\exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{U_{i,t}(\lambda_i)}(y)$, i.e., $\forall \lambda_{-i}$, $(u_{i,t}(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Theorem

If C is closed and convex and a B -set for player i for every t , then it is approachable by player i .

Approachability with Time-Varying Payoffs

Proof

Step 1: Projection Rule.

(1) Let $y_{t-1} \in P_C(\bar{u}_{i,t-1})$.

(2) If $\bar{u}_{i,t-1} \in C$, then play any $\lambda_{i,t} \in \Delta(A_i)$.

(3) If $\bar{u}_{i,t-1} \notin C$, note: C convex and closed $\implies (\bar{u}_{i,t-1} - y_{t-1}) \in N_C(y_{t-1}) \iff (\bar{u}_{i,t-1} - y_{t-1}) \cdot (z - y_{t-1}) \leq 0 \iff C \subseteq H^-(\bar{u}_{i,t-1}, y_{t-1})$.

(4) C convex, closed, B -set at $t+C \subseteq H^-(\bar{u}_{i,t-1}, y_{t-1}) \implies \exists \lambda_i \in \Delta(A_i) : \forall \lambda_{-i}, u_{i,t}(\lambda_i, \lambda_{-i}) \in \subseteq H^-(\bar{u}_{i,t-1}, y_{t-1}) \iff \lambda_{i,t} \in \Delta(A_i) : \forall \lambda_{-i} (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_{i,t}(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \leq 0$.

(5) Play $\lambda_{i,t}$ at stage t .

Steps 2 and 3 as before, just replacing u_i with $u_{i,t}$.

Approachability with Time-Varying Payoffs

Definition

C is strong **B-set** for player i if, $\forall t, \forall x \in \bigcup_t \text{co}(u_{i,t}) \setminus C, \exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{\bigcup_{i,t}(\lambda_i)}(y)$, i.e., $\forall \lambda_{-i}, (u_{i,t}(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Definition

C is **B-set** for player i at time t if, $\forall x \in \text{co}(u_{i,t}) \setminus C, \exists y \in P_C(x)$ and $\lambda_i \in \Delta(A_i)$ s.t. $(x - y) \in N_{\bigcup_{i,t}(\lambda_i)}(y)$, i.e., $\forall \lambda_{-i}, (u_{i,t}(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

With time-varying payoffs, it matters whether $\forall x \in \bigcup_t \text{co}(u_{i,t}) \setminus C$ or $\forall x \in \text{co}(u_{i,t}) \setminus C$!

Strong B-set (non-canonical terminology) + closed is enough for approachability.

B-set at every t + closed is NOT enough for approachability; counterexamples exist.

Need convexity to make it work in general.

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Approachability with Time-Varying Payoffs
- **Application: Picking Experts**
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Application: Picking Experts

S states of nature. A actions. Payoffs $u : A \times S \rightarrow \mathbb{R}$. Set of experts E .

Every period t ,

- (1) state s^t realises,
- (2) each expert recommends action $a_{e,t} \in A$,
- (3) DM chooses which expert to follow $e_t \in E$ and adopts their recommended action,
- (4) payoffs realise, and DM observes s_t .

Application: Picking Experts

A actions. S states of nature. Payoffs $u : A \times S \rightarrow \mathbb{R}$. Set of experts E .

History: $h_t \in H_t := (S \times A^{|E|} \times E)^{t-1}$ (previous states, what each expert recommended, expert chosen). $H := \cup_t H_t$.

Strategy: $\sigma : H \rightarrow \Delta(E)$.

Average payoff: $\bar{u}_T(\sigma) := \frac{1}{T} \sum_{t \leq T} \sum_e \sigma(h_t)(e) u(a_{e,t}, s_t)$.

Payoff from following particular expert e : $\bar{u}_T(e)$.

Small problem: DM doesn't know what experts actually know, whether have full info, partial, no info, biased, etc.

Definition

DM's σ is **no-regret strategy** if $\forall e \in E$ and each sequence $s_1, s_2, \dots$,

$$\mathbb{P}^\sigma \left(\liminf_{t \rightarrow \infty} \bar{u}_t(\sigma) - \bar{u}_t(e) \geq 0 \right) = 1.$$

Does no-regret strategy even exist? Can we characterise it?

Application: Picking Experts

Theorem

The DM has a no-regret strategy.

Proof

Opponent: nature, choosing $\sigma_0 : H \rightarrow \Delta(S)$. $C := \mathbb{R}_+^{|E|}$.

Let $u_t : E \times S \rightarrow \mathbb{R}$ be s.t. $u_t(e, s) = u(a_{e,t}, s)$.

For $\lambda \in \Delta(E)$ and $\gamma \in \Delta(S)$, let $v_t(\lambda, \gamma) := (u_t(\lambda, \gamma) - u_t(e, \gamma))_{e \in E} \in \mathbb{R}^{|E|}$.

Let $\bar{v}_T := \frac{1}{T} \sum_{t \leq T} v_t(\sigma(h_t), \sigma_0(h_t))$. *Regret vector*: no-regret $\iff \liminf_t \bar{v}_t \in C$.

For $x \in \mathbb{R}^{|E|}$, projection onto C is $y := x^+$ (positive part), and the normal is $x^- := y - x$ (negative part).

Choose the *Blackwell action* at x : if $\sum_e x_e^- > 0$, set

$$\lambda^x(e) := \frac{x_e^-}{\sum_{e'} x_{e'}^-} \quad (\text{put weight on experts relative to which you are behind}),$$

and any λ^x if $x \in C$. Note: λ^x = regret matching!

Application: Picking Experts – Proof (via Blackwell)

Proof

$$C := \mathbb{R}_+^{|E|}. \quad v_t(\lambda, \gamma) := (u_t(\lambda, \gamma) - u_t(e, \gamma))_{e \in E} \in \mathbb{R}^{|E|}; \quad \bar{v}_T := \frac{1}{T} \sum_{t \leq T} v_t(\sigma(h_t), \sigma_0(h_t)).$$

$$\text{For } x \in \mathbb{R}^{|E|}, y := x^+, x^- := y - x. \text{ For } \neg(x \geq 0), \text{ set } \lambda^x(e) := \frac{x_e^-}{\sum_{e'} x_{e'}^-}.$$

For any opponent choice $\gamma \in \Delta(S)$,

$$(i) \neg(x \geq 0) \implies \|x^-\| = \|x - y\| > 0;$$

$$(ii) 0 \geq (v_t(\lambda^x, \gamma) - y) \cdot (x - y) = (x^+ - v_t(\lambda^x, \gamma)) \cdot x^- = x^+ \cdot x^- - v_t(\lambda^x, \gamma) \cdot x^- = -v_t(\lambda^x, \gamma) \cdot x^-.$$

Note that

$$\begin{aligned} x^- \cdot v_t(\lambda^x, \gamma) &= \sum_e x_e^- \left[\sum_{e'} \lambda^x(e') u_t(e', \gamma) - u_t(e, \gamma) \right] \\ &= \sum_{e'} \frac{\sum_e x_e^-}{\sum_{e''} x_{e''}^-} x_{e'}^- u_t(e', \gamma) - \sum_e x_e^- u_t(e, \gamma) = \sum_{e'} x_{e'}^- u_t(e', \gamma) - \sum_e x_e^- u_t(e, \gamma) = 0. \end{aligned}$$

Hence C is closed, convex, B -set at every t . Blackwell's approachability theorem with time-varying payoffs $\implies \bar{v}_t$ approaches C a.s., i.e.,

$$\liminf_{t \rightarrow \infty} (\bar{u}_t(\sigma, \sigma_0) - \bar{u}_t(e, \sigma_0)) \geq 0 \text{ for all } e \in E \text{ and any of nature's moves } \sigma_0.$$

Application: Picking Experts

Theorem

The DM has a no-regret strategy.

Strategy (implementable): compute current average regrets $x := \bar{v}_{t-1}$; if $x \notin C$ play λ^x .

No need to know anything about the experts.

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Approachability with Time-Varying Payoffs
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

Hannan (1957) Consistency

Setup:

Two players face repeated game.

Actions: A_i ; **Stage-Game Payoffs:** $u_1 : A_1 \times A_2 \rightarrow \mathbb{R}$ $u_2 := -u_1$;

Histories: $H_t := A^t$, $H := \cup_t H_t$; **Strategies:** $\sigma_i : H \rightarrow \Delta(A_i)$; $\lambda_i \in \Delta(A_i)$.

Expected Payoffs: $u_i(\sigma(h_t))$.

Average Payoffs: $\bar{u}_{i,T}(\sigma) := \frac{1}{T} \sum_{t \leq T} u_i(a_t)$.

Average Expected Payoff: $\bar{u}_{i,T}(\sigma) := \frac{1}{T} \sum_{t \leq T} \sum_a \sigma(h_t)(a) u_i(a)$.

Benchmark: (External) Regret

External regret: $\max_{a_i \in A_i} \bar{u}_{i,t}(a_i) - \bar{u}_{i,t}(\sigma)$.

External regret: comparison relative to swapping to fixed action.

Hannan consistency: $\limsup_{T \rightarrow \infty} \max_{a_i \in A_i} \bar{u}_{i,t}(a_i) - \bar{u}_{i,t}(\sigma) \leq 0$ a.s.

Goal: show existence of Hannan consistent strategy.

Boils down to picking experts when each actions is consistently recommended by same expert.

Alternative via SFP

Smoothed Fictitious Play (SFP)

Let $U \in \mathbb{R}^A$, $c(\lambda) := D_{KL}(\lambda \| (1/|A|)) = \sum_a \lambda(a) \ln(\lambda(a)) + \ln(|A|)$, $V(U, \eta) := \max_{\lambda \in \Delta(A)} \lambda \cdot U - \eta c(\lambda)$, and $\lambda^*(U, \eta) := \arg \max_{\lambda \in \Delta(A)} \lambda \cdot U - \eta c(\lambda)$. For any $t \geq 2$, let $\sigma(h_t) \equiv \lambda_t := \lambda^*(U_{t-1}, \eta_t)$ where $U_t := (u(a, \tilde{\gamma}_t))_{a \in A}$ and $(\eta_t) : \eta_t \downarrow 0$.

Can generalise to additive perturbed utility with infinite mg cost at boundary of simplex.
Interpretation: best respond to slightly perturbed empirical model; perturbations vanish.
No ties.

Proposition 4.5 (Fudenberg and Levine (1998; 1999 GEB))

Under SFP, $\limsup_{t \rightarrow \infty} \max_{a_i \in A_i} \bar{u}_{i,t}(a_i) - \bar{u}_{i,t}(\sigma) \leq 0$ a.s.

Can also accommodate time-varying payoffs.

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Classical Calibration
- Calibrating Forecasts
- More on Calibration
- Calibration and Learning in Games

4. Adaptive Algorithms

Calibration

Important element of learning in games: forecasting opponent.

Calibration: from learning literature (Dawid 1982 JASA)

Suppose that, in a long conceptually infinite sequence of weather forecasts, we look at all those days for which the forecast probability of precipitation was, say, close to some given value p and assuming these form an infinite sequence determine the long run proportion $\mathbf{p}$ of such days on which the forecast event rain in fact occurred. The plot of $\mathbf{p}$ against p is termed the forecaster's empirical calibration curve. If the curve is the diagonal $\mathbf{p} = p$, the forecaster may be termed well calibrated.

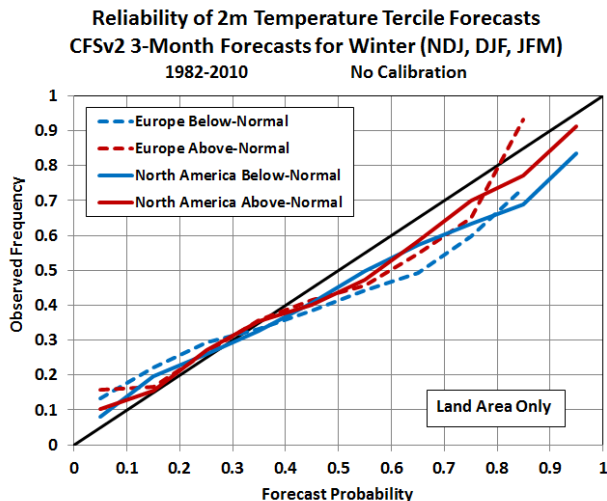
Not being calibrated is bad.

Dawid (1982 JASA): If data generated by probabilistic model, then forecasts generated by that model are a.s. calibrated.

Not being calibrated $\implies$ Have wrong probabilistic model.

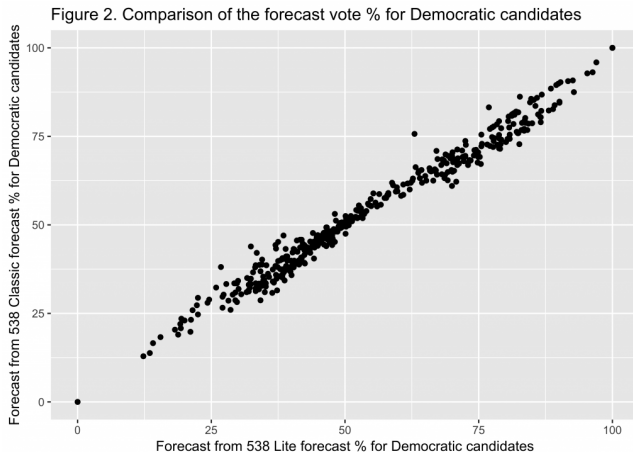
Statisticians, pollsters, forecasters want to be calibrated.

Aspiring to Calibration



<https://www.worldclimateservice.com/2020/07/06/what-is-forecast-reliability/>

Aspiring to Calibration



[https://www.science.org/content/blog-post/
analysis-prediction-results-united-states-congressional-elections-2018](https://www.science.org/content/blog-post/analysis-prediction-results-united-states-congressional-elections-2018)

Aspiring to Calibration

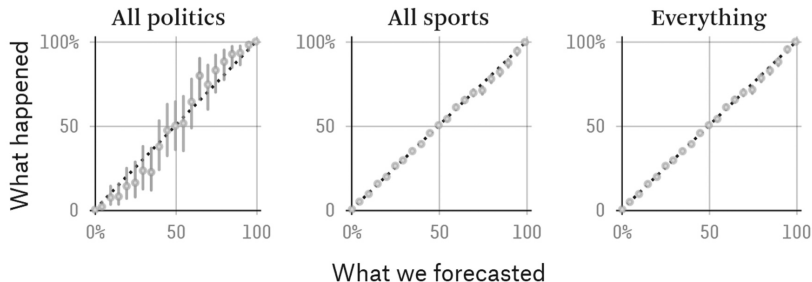


FIG. 2.—Calibration plots of FiveThirtyEight (projects.fivethirtyeight.com/checking-our-work, updated June 26, 2019). For example, in the “Everything” plot, the 10% data point (which lies slightly below the diagonal) has the following attached description: “We thought the 107,962 observations in this bin had a 10% chance of happening. They happened 9% of the time.” A color version of this figure is available online.

(Foster Hart 2021 JPE)

Aspiring to Calibration

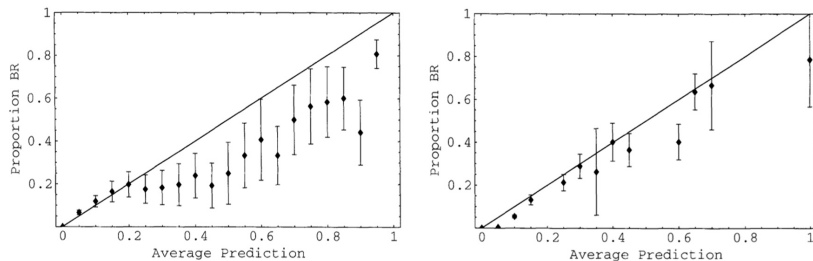


FIG. 1.—In Foster and Stine (2004), the business problem was to forecast the chance of a person going bankrupt in the next month. Both of the above forecasts are based on a large linear model. The one on the left was obtained by a logistic regression and the one on the right by a monotone regression. The left-hand forecast is not calibrated, whereas the right-hand forecast is calibrated and so can be used directly for decision-making. BR = bankruptcy.

(Foster Hart 2021 JPE)

Aspiring to Calibration

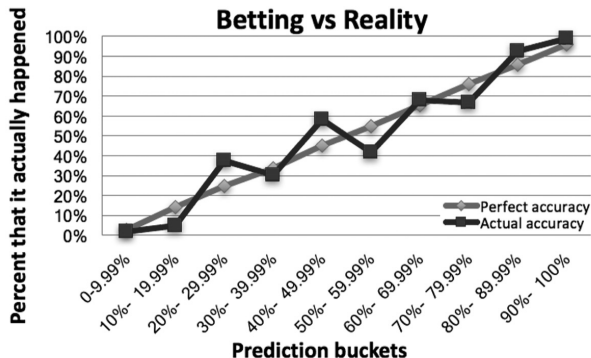
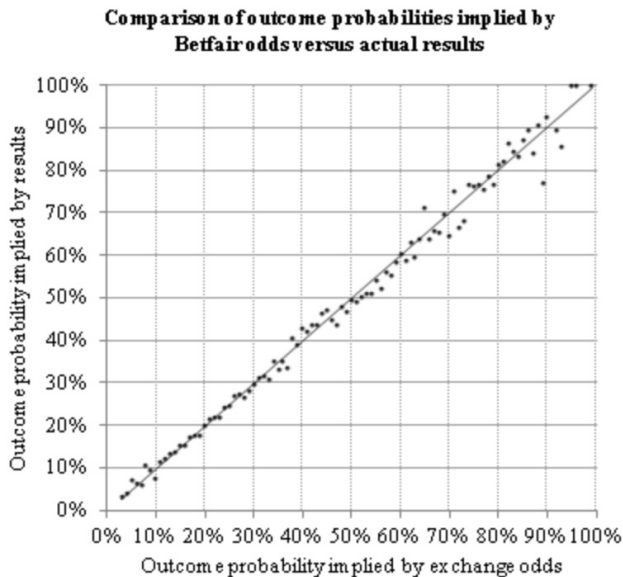


FIG. 3.—Calibration plot of ElectionBettingOdds (electionbettingodds.com/TrackRecord.html, updated November 13, 2018), which “tracked some 462 different candidate chances across dozens of races and states in 2016 and 2018.” A color version of this figure is available online.

(Foster Hart 2021 JPE)

Aspiring to Calibration



https://www.football-data.co.uk/blog/wisdom_of_the_crowd.php

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Classical Calibration
- Calibrating Forecasts
- More on Calibration
- Calibration and Learning in Games

4. Adaptive Algorithms

Calibrated Learning: Setup

Stage Game

Players $i \in \{1, 2\}$; actions A_i finite; $A = A_1 \times A_2$.

Payoffs $u_i : A \rightarrow \mathbb{R}$.

Forecasts: $\Delta(A_{-i})$.

Repeated Play: $t = 1, 2, \dots$

History $H_t := (\Delta(A_{-i}) \times A)^{t-1}$, $H := \cup_t H_t$.

Strategies $\sigma_i : H \rightarrow \Delta(A_i)$; realised actions $a_t = (a_{1,t}, a_{2,t})$.

Empirical distribution $\bar{\sigma}_t \in \Delta(A)$: $\bar{\sigma}_t(a) := \frac{1}{t} \sum_{s \leq t} \mathbf{1}_{\{a_s = a\}}$.

Beliefs and Behaviour

Player i 's beliefs $\sigma_{-i}^i : H \rightarrow \Delta(A_{-i})$.

Myopic best replies: $a_{i,t} \in \arg \max_{a'_i} u_i(a'_i, \sigma_{-i}^i)$

Fix a deterministic tie-breaking rule.

Forecasting Rule: $\phi_i : H \rightarrow F$.

Deterministic: $\phi_i : H \rightarrow F \subseteq \Delta(A_{-i})$,

i.e. $f_{i,t} \equiv \phi_i(h_t) = \sigma_{-i,t}^i \in F$.

Stochastic: $\phi_i : H \rightarrow \Delta(F)$. Realised forecast $f_{i,t} \sim \phi_i(h_t)$.

Calibration

Assume F finite; $F = \{\sigma_{-i}^1, \dots, \sigma_{-i}^M\}$.

Forecast Counts: $n_{i,t}^m := \sum_{s \leq t} \mathbf{1}_{\{f_{i,s} = \sigma_{-i}^m\}}$.

Conditional Empirical Frequency: $\bar{\sigma}_{-i,t}^m := \frac{\sum_{s \leq t} a_{-i,s} \mathbf{1}_{\{f_{i,s} = \sigma_{-i}^m\}}}{n_{i,t}^m}$ whenever $n_{i,t}^m > 0$.

Calibration Error of Forecast m : $k_{i,t}^m := \|\bar{\sigma}_{-i,t}^m - \sigma_{-i}^m\| n_{i,t}^m / t$ if $n_{i,t}^m > 0$; $k_{i,t}^m := 0$ ow.

Calibration Score: $K_{i,t} := \sum_m k_{i,t}^m$.

Definition (Foster and Vohra 1998 Biometrika)

Forecasting rule ϕ_i is **ε -calibrated** if for every σ_{-i}

$$\limsup_{t \rightarrow \infty} K_{i,t} < \varepsilon \text{ a.s.}$$

Remark (Oakes 1985 JASA; Dawid 1985 JASA)

For small enough ϵ , no deterministic forecasting rule can be ϵ -calibrated.

Theorem (Foster and Vohra 1998 Biometrika)

$\forall \epsilon > 0$, there is ϵ -calibrated forecasting rule.

Calibrated Forecasts

Foster recalls that “this paper took the longest to get published of any I have worked on. I think our first submission was about 1991. Referees simply did not believe the theorem – so they looked for amazingly tiny holes in the proof. When the proof had been compressed from its original 15-20 pages down to about 1, it was finally believed.” (in Olszewski 2015)

Proving the calibration rule theorem

Original proof: Foster and Vohra (1998 Biometrika).

Hart (2023): proof via minmax theorem (existence).

Foster (1999 GEB): proof via Blackwell approachability.

Fudenberg and Levine (1999 GEB): specific construction.

Proof Idea

Quadratic Cost of Forecast m : $c_{i,t}^m := \|\bar{\sigma}_{-i,t}^m - \sigma_{-i}^m\|^2 n_{i,t}^m$ if $n_{i,t}^m > 0$; $c_{i,t}^m := 0$ ow.

Total Cost: $C_{i,t} := \sum_m c_{i,t}^m$.

Fix a partition Π_i of $\Delta(A_{-i})$ where:

$\Pi_i = \{B_{-i}^m\}_{m=1}^M$ of $\Delta(A_{-i})$, for some finite M ;

Each B_{-i}^m is convex and $\sigma_{-i}^m \in B_{-i}^m$ is the centroid of B_{-i}^m .

Fineness of grid: $\max_m \max_{\sigma_{-i} \in B_{-i}^m} \|\sigma_{-i}^m - \sigma_{-i}\| \leq \eta$.

Write $\Pi_i(M, \eta)$

Player i chooses $\phi_i : H \rightarrow F$.

Proof Idea

Quadratic Cost of Forecast m : $c_{i,t}^m := \|\bar{\sigma}_{-i,t}^m - \sigma_{-i}^m\|^2 n_{i,t}^m$ if $n_{i,t}^m > 0$; $C_{i,t} := \sum_m c_{i,t}^m$.

$$g_t(m, a_{-i}) := \mathbb{E}[C_{i,t} - C_{i,t-1} \mid H_{t-1}, f_{i,t} = \sigma_{-i}^m, a_{-i,t} = a_{-i}] = (n_{i,t-1}^m + 1) \left\| \frac{\bar{\sigma}_{-i,t}^m + 1_m}{n_{i,t-1}^m + 1} - \sigma_{-i}^m \right\|^2 - c_{i,t-1}^m.$$

Let g_t be payoffs in zero-sum game where i chooses $\lambda \in \Delta(M)$ and nature $\gamma \in \Delta(A_{-i})$.

T -initialised myopic strategies:

- (1) Initialisation Phase: Each pure strategy repeated T times.
- (2) At $t > TM$, choose $\lambda_t = \text{maxmin strategy}$ in zero-sum game with payoffs g_t .

Theorem (Fudenberg and Levine (1999 GEB))

For any ε , $\exists \Pi_i(M, \eta)$ s.t. T -initial myopic strategy satisfies $\limsup C_{i,t}/t \leq \varepsilon$ a.s.

Lemma (Fudenberg and Levine (1999 GEB))

After the initialisation phase, a T -initial myopic strategy has $g_t \leq \frac{1}{T+1} + \eta^2$.

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Classical Calibration
- Calibrating Forecasts
- More on Calibration
- Calibration and Learning in Games

4. Adaptive Algorithms

Calibration

Day	1	2	3	4	5	6	...	$\mathcal{K}$	$\mathcal{R}$	$\mathcal{B}$
Rain	1	0	1	0	1	0				
F1	100%	0%	100%	0%	100%	0%		0	0	0
F2	50%	50%	50%	50%	50%	50%		0	0.25	0.25

FIGURE 1. Two calibrated forecasts.

(Foster Hart 2023 TE)

Calibration Score: $\mathcal{K}_t := \frac{1}{t} \sum_{\ell \leq t} \sum_m \mathbf{1}_{\{f_\ell = \sigma^m\}} \|\bar{\sigma}_\ell^m - f_\ell\|^2$.

Calibration: getting $\mathcal{K}_t$ arbitrarily close to 0.

Two calibrated forecasts can be wildly different regarding accuracy of predictions.

Quadratic Scoring Rule

Day	1	2	3	4	5	6	...	$\mathcal{K}$	$\mathcal{R}$	$\mathcal{B}$
Rain	1	0	1	0	1	0				
F1	100%	0%	100%	0%	100%	0%		0	0	0
F2	50%	50%	50%	50%	50%	50%		0	0.25	0.25

FIGURE 1. Two calibrated forecasts.

Brier Score (1950): $\mathcal{B}_t := \frac{1}{t} \sum_{\ell \leq t} \|a_\ell - f_\ell\|^2$.

Original QSR of belief elicitation. F1 has $\mathcal{B}_t = 0$; F2 has $\mathcal{B}_t = 1/4$.

Refinement Score: $\mathcal{R}_t := \frac{1}{t} \sum_{\ell \leq t} \sum_m \mathbf{1}_{\{f_\ell = \sigma^m\}} \|\bar{\sigma}_\ell^m - a_\ell\|^2$.

Captures within-bin variance of forecasts.

$\mathcal{B}_t = \mathcal{R}_t + \mathcal{K}_t$. ($\mathbb{E}[X^2] = \mathbb{V}(X) + \mathbb{E}[X]^2$).

Is there natural trade-off between calibration and refinement?

For any forecast rule, is it possible to calibrate $\mathcal{K}_t$ without increasing $\mathcal{R}_t$?

Offline, yes. Take forecast rule and, for each bin m , readjust the forecast σ^m to correspond to within-bin historical frequency. Keep $\mathcal{R}_t$ (within-bin variance) and eliminate $\mathcal{K}_t$.

‘Calibeating’: Reducing Brier score by amount equal to calibration score.

And procedures to do this online, i.e, on-the-fly?

Foster Hart (2023 TE), Calibeating Beating forecasters at their own game

Typical calibrated algorithms: get zero calibration score, but do rather poorly at refinement score. Would suggest trade-off.

Point of the paper: for any forecast rule, it is possible to obtain a calibrated forecast ($\mathcal{K}_t \rightarrow 0$) without increasing the refinement score.

Procedure 1: replace each forecast by empirical frequency on previous days in which this forecast was made.

Issue: Procedure may not yield calibrated forecasts. Can do better.

Procedure 2: Self-calibeating = Calibrating; issue of infinite regress.

Calibeating by calibrated forecast via stochastic forecast-hedging calibration for each bin.

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Classical Calibration
- Calibrating Forecasts
- **More on Calibration**
- Calibration and Learning in Games

4. Adaptive Algorithms

Detecting the Expert

Issue: calibration seems to suggest we cannot differentiate between expert and ignorant.

Not exactly true. E.g., get putative expert to submit theory at time 0; can find tests that cannot be ignorantly passed on almost all data sets (Olszewski and Sandroni 2009 AnnStat, A nonmanipulable test).

Bayesian tester approach: Stewart 2011 JET, Nonmanipulable Bayesian testing.

Relation Between Calibration and Learning

Kalai, Lehrer, and Smorodinsky (1999 GEB): equivalence between different notions of merging and of calibration. (Will make more sense later on.)

Olszewski (2015 Ch), Calibration and Expert Testing: comprehensive survey on calibration.

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Classical Calibration
- Calibrating Forecasts
- More on Calibration
- Calibration and Learning in Games

4. Adaptive Algorithms

Convergence issues of the learning in games:

Generalised FP: *if* converge, asymptotic behaviour is Nash-like; but convergence not assured.

Similar issues with replicator dynamic and other models.

Foster and Vohra (1997 GEB): different learning basis – *calibration* – yields different solution concept – correlated equilibrium.

Correlated Equilibrium

Players I . Actions A_i ; $A = \times_i A_i$. Stage Payoffs $u_i : A \rightarrow \mathbb{R}$.

Correlated Equilibrium: $p \in \Delta(A) : \sum_{a_{-i}} p(a_i, a_{-i})(u_i(a_i, a_{-i}) - u_i(a'_i, a_{-i})) \geq 0,$
 $\forall i \in I, \forall a_i, a'_i \in A_i.$

Correlated ε -Equilibrium: $p \in \Delta(A) : \sum_{a_{-i}} p(a_i, a_{-i})(u_i(a_i, a_{-i}) - u_i(a'_i, a_{-i})) \geq -\varepsilon,$
 $\forall i \in I, \forall a_i, a'_i \in A_i.$

Game repeated $t = 1, 2, \dots$. Stage payoffs $u_i(a_{i,t}, a_{-i,t})$.

Empirical frequency: $\bar{p}_t \in \Delta(A)$ s.t. $\bar{p}_t(a) = \frac{1}{t} \sum_{\ell \leq t} \mathbf{1}_{\{a_\ell = a\}}.$

Theorem 1 (Foster and Vohra 1997 GEB)

Suppose each player uses a calibrated forecasting scheme (on arbitrarily fine partitions) and in each t plays myopic best reply to their forecast (fixed tie-breaking). Then empirical frequency of play $\bar{p}_t \in \Delta(A)$ converges to set of correlated equilibria of stage game.

Idea

Take player i 's forecasts as signal.

In the limit, conditional on forecast, joint distribution of opponents coincides with forecast.

Best-response to forecast means no profitable deviation conditional on signal.

Calibrated Learning $\implies$ Correlated Equilibrium

Theorem 1 (Foster and Vohra 1997 GEB)

Suppose each player uses a calibrated forecasting scheme (on arbitrarily fine partitions) and in each t plays myopic best reply to their forecast (fixed tie-breaking). Then empirical frequency of play $\bar{p}_t \in \Delta(A)$ converges to set of correlated equilibria of stage game.

Meaning of calibration: whenever you forecast p , reality looks like p on that subsequence.

Behavioural content: minimal discipline on beliefs + myopic optimality $\implies$ CE.

Internal vs external regret: no internal regret also leads to CE.

Design/selection: by designing calibrated grids, any target σ can be attained.

Attainability of Any Correlated Equilibrium

Theorem 2 (Foster and Vohra 1997 GEB)

For any correlated equilibrium, there exist calibrated forecasts and myopic best replies such that empirical frequency of play $\bar{p}_t \in \Delta(A)$ converges to that correlated equilibrium.

Idea

Construct partitions and representatives matching correlated equilibrium conditionals; calibrate to them.

Best replies implement the recommended supports; empirical play tracks correlated equilibrium.

Overview

1. Learning in Games
2. Approachability
3. Calibration
4. Adaptive Algorithms

Learning Dynamics

Actions: $A_i, A := \times_i A_i$. $u_i : A \rightarrow \mathbb{R}$. I players. $X \subseteq \mathbb{R}^m$ convex; $X = \times_i A_i$ or $X = \times_i \Delta(A_i)$.

Learning dynamic: $\dot{x}_t = F(x_t; u)$; $\dot{x}_i = F_i(x; u)$; $F \in \mathcal{C}^1$.

Uncoupled dynamics: $\dot{x}_i$ only depends on x and u_i , not on others' payoffs; $\dot{x}_i = F_i(x; u_i)$.

$\mathcal{U} \subseteq \mathbb{R}^A$ has **unique NE property** if $\forall u \in \mathcal{U}$, $\langle I, A, u \rangle$ has unique NE.

F is **Nash-convergent** for $\mathcal{U}$ if $\forall u \in \mathcal{U}$, the learning dynamics always converges to NE, i.e., $F(x^*; u) = 0$ for any NE x^* and $\lim_{t \rightarrow \infty} x_t = x^*$ for any x_0 .

Theorem (Hart and Mas-Colell 2003 AER)

$\exists u_0$ such that for any $\mathcal{U}$ containing a neighbourhood of u_0 , no uncoupled dynamics is Nash-convergent for $\mathcal{U}$. Furthermore, there is $\mathcal{U}$ containing a neighbourhood of u_0 s.t. $\mathcal{U}$ has unique NE property.

E.g., u_0 is game for which FP doesn't work (Jordan 1993 GEB).

Theorem (Hart and Mas-Colell 2003 AER)

$\exists u_0$ such that for any $\mathcal{U}$ containing a neighbourhood of u_0 , no uncoupled dynamics is Nash-convergent for $\mathcal{U}$. Furthermore, there is $\mathcal{U}$ containing a neighbourhood of u_0 s.t. $\mathcal{U}$ has unique NE property.

- $\exists$ uncoupled dynamics guaranteeing NE convergence for some families of games (e.g. SFP for zero-sum, potential games, supermodular games, etc).
- $\exists$ uncoupled dynamics that are most of the time close to NE, but not Nash-convergent: exit infinitely often any neighborhood of NE (Foster and Young 2003 GEB).
 - Smooth calibrated learning dynamics gets $(1 - \epsilon)$ -close to ϵ -NE. (Foster and Hart 2018 GEB, Theorem 15)
 - Calibration + ϵ -BR gets ϵ -NE. (Foster and Hart 2021 JPE, Theorem 13)
 - Both are uncoupled dynamics.

Learning (Correlated) Equilibria (bis)

Correlated Equilibria:

Nash equilibrium of game with signals.

Epistemic foundations (Aumann 1974 JMathE; 1987 Ecta).

Available correlating signals may simply make their way into strategic behaviour.

Hart and Mas-Colell (2000 Ecta): simple learning procedure that converges to CE a.s.

Correlated Equilibrium

Players I . Actions A_i ; $A = \times_i A_i$. Stage Payoffs $u_i : A \rightarrow \mathbb{R}$.

Correlated Equilibrium: $p \in \Delta(A) : \sum_{a_{-i}} p(a_i, a_{-i})(u_i(a_i, a_{-i}) - u_i(a'_i, a_{-i})) \geq 0,$
 $\forall i \in I, \forall a_i, a'_i \in A_i.$

Correlated ε -Equilibrium: $p \in \Delta(A) : \sum_{a_{-i}} p(a_i, a_{-i})(u_i(a_i, a_{-i}) - u_i(a'_i, a_{-i})) \geq -\varepsilon,$
 $\forall i \in I, \forall a_i, a'_i \in A_i.$

Game repeated $t = 1, 2, \dots$. Stage payoffs $u_i(a_{i,t}, a_{-i,t})$.

Empirical frequency: $\bar{p}_t \in \Delta(A)$ s.t. $\bar{p}_t(a) = \frac{1}{t} \sum_{\ell \leq t} \mathbf{1}_{\{a_\ell = a\}}.$

Regret Matching

At time t consider counterfactual play: replacing past play of $\hat{a}_i$ with a_i .

$w_{i,t}(\hat{a}_i, a_i) := u_i(a_i, a_{-i,t})$ if $a_{i,t} = \hat{a}_i$; otherwise $w_{i,t}(\hat{a}_i, a_i) := u_i(a'_i, a_{-i,t})$ if $a_{i,t} \neq \hat{a}_i$.

Average Payoff Difference: $d_{i,t}(\hat{a}_i, a_i) := \frac{1}{t} \sum_{\ell \leq t} [w_{i,\ell}(\hat{a}_i, a_i) - u_i(a_\ell)]$.

Average Regret: $r_{i,t}(\hat{a}_i, a_i) := d_{i,t}(\hat{a}_i, a_i)^+$.

Regret Matching: Adjust behaviour more toward actions that regret not having taken:
Switch next period to different action with probability proportional to regret for that action, i.e., increase in payoff had such change always been made in the past.

$$\lambda_{i,t+1}(a_i) := \frac{1}{\eta} r_{i,t}(a_{i,t}, a_i) \quad \forall a_i \neq a_{i,t}; \quad \text{and} \quad \lambda_{i,t+1}(a_{i,t}) := 1 - \sum_{a_i \neq a_{i,t}} \lambda_{i,t+1}(a_i);$$
$$\eta > 2 \|u_i\|_\infty |A_i|.$$

η : measure of inertia; higher $\eta \implies$ lower switching probability.

Very behavioural strategy.

Theorem (Hart and Mas-Colell 2000 Ecta)

If all players play according to regret matching, empirical play frequencies converge to set of CE of stage game a.s.

Regret Matching Implies CE

Remark

Converge to regret below ϵ iff converge to ϵ -CE.

Proof

For converging $\bar{p}_t$, $\limsup_{t \rightarrow \infty} r_{i,t}(\hat{a}_i, a_i) \leq \epsilon \forall i, \hat{a}_i, a_i$, if and only if
 $\limsup_{t \rightarrow \infty} d_{i,t}(\hat{a}_i, a_i) = \sum_{a_{-i}} \bar{p}_t(\hat{a}_i, a_{-i})(u_i(a_i, a_{-i}) - u_i(\hat{a}_i, a_{-i})) \leq \epsilon, \forall i, \hat{a}_i, a_i.$

Theorem (Hart and Mas-Colell 2000 Ecta)

If all players play according to regret matching, empirical play frequencies converge to set of CE of stage game a.s.

Proof via straight up applying Blackwell's approachability with payoff vector

$$v_i(a)(\tilde{a}_i, \hat{a}_i) = (\mathbf{1}_{\{a_i = \hat{a}_i\}}[u_i(\tilde{a}_i, a_{-i}) - u_i(\hat{a}_i, a_{-i})])_{\tilde{a}_i, \hat{a}_i \in A_i} \in \mathbb{R}^{|A_i|^2}$$

Speed of convergence: $\mathbb{E}[r_{i,t}(\hat{a}_i, a_i)] \leq Kt^{-1/2}$ ($O(t^{-1/2})$).

Can't do better with stationary mixed strategies (errors of order $t^{-1/2}$ by CLT).

$\mathbb{P}(\bar{p}_t \text{ is } \varepsilon - \text{CE}, \forall t > T) \geq 1 - \delta \exp(-\delta T)$, where $\delta = \delta(\varepsilon)$. (Large deviation result)

Theorem (Hart and Mas-Colell 2000 Ecta)

If all players play according to regret matching, empirical play frequencies converge to set of CE of stage game a.s.

Can extend result to case in which game is not known and others' choices not observed. Only observe $u_{i,t}(= u_i(a_{i,t}, a_{-i,t}))$.

Replace $d_{i,t}(\hat{a}_i, a_i)$ with $\tilde{d}_{i,t}(\hat{a}_i, a_i) := \frac{1}{t} \left[\sum_{\ell \leq t: a_{i,\ell} = a_i} \frac{\bar{p}_{i,\ell}(\hat{a}_i)}{\bar{p}_{i,\ell}(a_i)} u_{i,\ell} \right] - \frac{1}{t} \left[\sum_{\ell \leq t: a_{i,\ell} = \hat{a}_i} u_{i,\ell} \right]$.

See Hart and Mas-Colell (2000 Ecta) for details.

If $\lambda_{i,t+1}(a_i) = f(r_{i,t}(a_{i,t}, a_i))$, for f Lipschitz continuous and sign-preserving ($x > (=) 0 \implies f(x) \geq (>) 0$), convergence to CE goes through (Cahn 2004 IJGT; Hart 2005 Ecta).

Theorem (Hart and Mas-Colell 2000 Ecta)

If all players play according to regret matching, empirical play frequencies converge to set of CE of stage game a.s.

It is thus interesting that Nash equilibrium, a notion that does not predicate coordinated behavior, cannot be guaranteed to be reached in an uncoupled way, while correlated equilibrium, a notion based on coordination, can. (Hart and Mas-Colell 2003 AER)

‘Conservation Coordination Law’ for game dynamics: some form of “coordination” must be present, either in the limit static equilibrium concept (such as correlated equilibrium) or in the dynamic leading to it (such as Nash equilibrium dynamics).

Approachability, Calibration, and Adaptive Algorithms

Duarte Gonçalves

University College London

Topics in Economic Theory