

Approachability, Calibration, Adaptive Algorithms, and Sophisticated Learning

Duarte Gonçalves

University College London

Topics in Economic Theory

Overview

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

Overview

1. Learning in Games

2. Approachability

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

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
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

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
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

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.

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 $[x] := \{y \in X | x \sim y\}$ and the set of equivalence classes $\hat{X} := \{[x], x \in X\}$.

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 $[x] := \{y \in X | x \sim y\}$ 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.

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.

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.

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.

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.

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

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.

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 $\hat{\succsim}$.

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 $\dim(X, \succsim) = 1$.

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.

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

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

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 $\forall x, y \in X, x \succsim y \iff u(x) \geq u(y)$.

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 $\dim(X, \succsim) < \infty$.

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 $x \succsim y \iff u(x) \geq u(y) \forall u \in U$.

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**
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

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:** $\sigma_i : H \rightarrow \Delta(A_i)$; $\lambda_i \in \Delta(A_i)$.

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

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

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. $\forall \epsilon > 0, \exists T : \forall \sigma_{-i}, \mathbb{P}^{\sigma}(d(\bar{u}_t, C) < \epsilon, \forall t \geq T) > 1 - \epsilon$; 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 .

Definition

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

approachable by player i if $\exists \sigma_i$ s.t. $\forall \epsilon > 0, \exists T : \forall \sigma_{-i}, \mathbb{P}^\sigma(d(\bar{u}_t, C) < \epsilon, \forall t \geq T) > 1 - \epsilon$; 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 wp1 uniformly over opponent's strategies: $\mathbb{P}^\sigma(\lim_{t \rightarrow \infty} d(\bar{u}_t, C) = 0) = 1$.

Definition

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

approachable by player i if $\exists \sigma_i$ s.t. $\forall \epsilon > 0, \exists T : \forall \sigma_{-i}, \mathbb{P}^\sigma(d(\bar{u}_t, C) < \epsilon, \forall t \geq T) > 1 - \epsilon$; 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 wp1 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.

Hyperplanes

Definition

Hyperplane $H(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x = b\}$.

Half-spaces: $H^+(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x \geq b\}$, $H^-(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x \leq b\}$.

$$H^+(a, b) \cap H^-(a, b) = H(a, b).$$

$$H^+(a, b) = H^-(-a, -b) \text{ and } H(a, b) = H(-a, -b).$$

Hyperplanes

Definition

Hyperplane $H(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x = b\}$.

Half-spaces: $H^+(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x \geq b\}$, $H^-(a, b) := \{x \in \mathbb{R}^m \mid a \cdot x \leq b\}$.

$$H^+(a, b) \cap H^-(a, b) = H(a, b).$$

$$H^+(a, b) = H^-(-a, -b) \text{ and } H(a, b) = H(-a, -b).$$

Definition

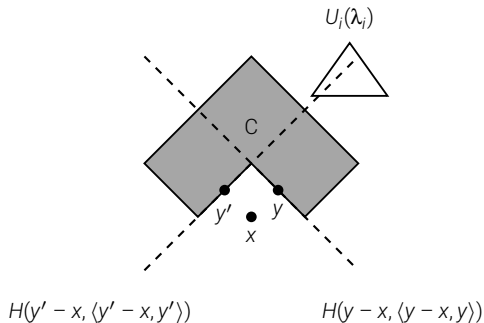
Hyperplane $H(a, b)$ separates x from C if (i) $x \in H^+(a, b) \setminus H(a, b)$ and $C \subseteq H^-(a, b)$, or (ii) $x \in H^-(a, b) \setminus H(a, b)$ and $C \subseteq H^+(a, b)$.

Remark

Hyperplane $H(x - y, (x - y) \cdot y)$:

- (i) $y \in H(x - y, (x - y) \cdot y)$.
- (ii) Orthogonal/Perpendicular to line passing through x and y , i.e. $z \in H(x - y, (x - y) \cdot y) \iff (z - y) \cdot (x - y) = 0$.
- (iii) y is point in $H(x - y, (x - y) \cdot y)$ closest to x : $\forall z \in H(x - y, (x - y) \cdot y), \|x - z\|^2 = \|x - y\|^2 + \|y - z\|^2$. (Pythagorean theorem)
- (iv) $\forall$ hyperplane H and $x \notin H$, if $y \in H$ is closest in H to x , then $H = H(x - y, (x - y) \cdot y)$.

Separating Hyperplanes in the B -set Condition



Hyperplane $H(y' - x, \langle y' - x, y' \rangle)$ separates x from $U_i(\lambda_i)$.

Definition

Closed set C is **B-set** for player i if $\forall x \in \text{co}(u_i) \setminus C, \exists y \in C$ and $\lambda_i \in \Delta(A_i)$ s.t.

(1) y is closest to x in C : $d(x, y) = d(x, C)$; and

(2) hyperplane $H(x - y, (x - y) \cdot y)$ separates x from $U_i(\lambda_i)$:

(i) $x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y)$ and (ii) $U_i(\lambda_i) \subseteq H^-(x - y, (x - y) \cdot y)$.

(ii) is equiv. to $\forall \lambda_{-i}, (u_i(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Definition

Closed set C is **B-set** for player i if $\forall x \in \text{co}(U_i) \setminus C, \exists y \in C$ and $\lambda_i \in \Delta(A_i)$ s.t.

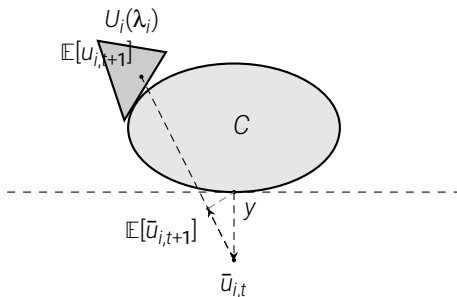
- (1) y is closest to x in C : $d(x, y) = d(x, C)$; and
- (2) hyperplane $H(x - y, (x - y) \cdot y)$ separates x from $U_i(\lambda_i)$:
 - (i) $x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y)$ and
 - (ii) $U_i(\lambda_i) \subseteq H^-(x - y, (x - y) \cdot y)$.

(ii) is equiv. to $\forall \lambda_{-i}, (u_i(\lambda_i, \lambda_{-i}) - y) \cdot (x - y) \leq 0$.

Theorem (Blackwell 1956)

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

Proving Approachability Theorem: One-Step Projection



Setup for One Step

- (1) Fix $t \geq 1$. Let $\bar{u}_{i,t-1} \in \text{co}(u_i)$.
- (2) If $\bar{u}_{i,t-1} \in C$, play anything (e.g., a_1). Otherwise, let $y_{t-1} \in C$ be a closest point:
 $d(\bar{u}_{i,t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.
- (3) By B -set, $\exists \lambda_{i,t} \in \Delta(A_i)$ s.t. hyperplane $H(\bar{u}_{i,t-1} - y_{t-1}, (\bar{u}_{i,t-1} - y_{t-1}) \cdot y_{t-1})$ separates $\bar{u}_{i,t-1}$ from $U_i(\lambda_{i,t})$:

$$(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \leq 0, \quad \forall \lambda_{-i} \in \Delta(A_{-i}).$$
- (4) Play $\lambda_{i,t}$ at stage t .

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Identity:

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

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Identity:

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

Write $(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - \bar{u}_{i,t-1}) = (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) - \|\bar{u}_{i,t-1} - y_{t-1}\|^2$.

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Identity:

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

Write $(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - \bar{u}_{i,t-1}) = (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) - \|\bar{u}_{i,t-1} - y_{t-1}\|^2$.

Conditional expectation: by separation,

$$\mathbb{E}[(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) \mid H_{t-1}] \leq \max_{\lambda_{-i}} (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \leq 0.$$

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Identity:

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

Write $(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - \bar{u}_{i,t-1}) = (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) - \|\bar{u}_{i,t-1} - y_{t-1}\|^2$.

Conditional expectation: by separation,

$$\mathbb{E}[(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) \mid H_{t-1}] \leq \max_{\lambda_{-i}} (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \leq 0.$$

Let $L := \max_{a \in A} \|u_i(a)\| < \infty$. Then $\mathbb{E}[\|u_i(a_t) - \bar{u}_{i,t-1}\|^2 \mid H_{t-1}] \leq 4L^2$.

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

Proving Approachability Theorem: One-Step Projection Inequality

Key Inequality

Update: $\bar{u}_{i,t} = \bar{u}_{i,t-1} + \frac{1}{t} (u_i(a_t) - \bar{u}_{i,t-1})$. Recall $d(\bar{u}_{t-1}, C) = \|\bar{u}_{i,t-1} - y_{t-1}\|$.

Identity:

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

Write $(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - \bar{u}_{i,t-1}) = (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) - \|\bar{u}_{i,t-1} - y_{t-1}\|^2$.

Conditional expectation: by separation,

$$\mathbb{E}[(\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(a_t) - y_{t-1}) \mid H_{t-1}] \leq \max_{\lambda_{-i}} (\bar{u}_{i,t-1} - y_{t-1}) \cdot (u_i(\lambda_{i,t}, \lambda_{-i}) - y_{t-1}) \leq 0.$$

Let $L := \max_{a \in A} \|u_i(a)\| < \infty$. Then $\mathbb{E}[\|u_i(a_t) - \bar{u}_{i,t-1}\|^2 \mid H_{t-1}] \leq 4L^2$.

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

Since $d(\bar{u}_{i,t}, C) \leq \|\bar{u}_{i,t} - y_{t-1}\|$, setting $V_t := d(\bar{u}_{i,t}, C)^2$ gives

$$\mathbb{E}[V_t \mid H_{t-1}] \leq \left(1 - \frac{2}{t}\right) V_{t-1} + \frac{4L^2}{t^2}.$$

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.

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.

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

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 $\mathbb{E}[V_t | \mathcal{H}_{t-1}] \leq (1 - \alpha_t)V_{t-1} + \beta_t$ 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.

Concluding the Proof of Blackwell's Approachability Theorem

- (1) Use the one-step choice $\lambda_{i,t}$ from the B -set condition at $\bar{u}_{i,t-1}$.
- (2) $\mathbb{E}[V_t \mid H_{t-1}] \leq (1 - \frac{2}{t})V_{t-1} + \frac{4L^2}{t^2}$.
- (3) Take $\alpha_t = \frac{2}{t}$ (diverges) and $\beta_t = \frac{4L^2}{t^2}$ (summable).
- (4) Robbins–Siegmund $\implies V_t = d(\bar{u}_t, C)^2 \rightarrow 0$ a.s.

Proving Approachability Theorem: Concluding

Concluding the Proof of Blackwell's Approachability Theorem

- (1) Use the one-step choice $\lambda_{i,t}$ from the B -set condition at $\bar{u}_{i,t-1}$.
- (2) $\mathbb{E}[V_t \mid H_{t-1}] \leq (1 - \frac{2}{t})V_{t-1} + \frac{4L^2}{t^2}$.
- (3) Take $\alpha_t = \frac{2}{t}$ (diverges) and $\beta_t = \frac{4L^2}{t^2}$ (summable).
- (4) Robbins–Siegmund $\implies V_t = d(\bar{u}_t, C)^2 \rightarrow 0$ a.s.

Theorem (Blackwell 1956)

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

Strategy (Blackwell's rule): at each t , project $\bar{u}_{i,t-1}$ onto C , pick $\lambda_{i,t}$ separating $\bar{u}_{i,t-1}$ from $U_i(\lambda_{i,t})$, play $\lambda_{i,t}$.

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 .

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Application: Picking Experts
- Universal Consistency

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

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

S states of nature. A actions. 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: $\sigma : H \rightarrow \Delta(E)$. Distribution of $s_t \sim \gamma_t \in \Delta(S)$.

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

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

Application: Picking Experts

S states of nature. A actions. 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: $\sigma : H \rightarrow \Delta(E)$. Distribution of $s_t \sim \gamma_t \in \Delta(S)$.

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

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

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 $s_1, s_2, \dots$,

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

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

Application: Picking Experts

Theorem

The DM has a no-regret strategy.

Simplifying assumption: $|E| = |A|$ and for each e , $a_{e,t} = a$ for some different a .

Application: Picking Experts

Theorem

The DM has a no-regret strategy.

Simplifying assumption: $|E| = |A|$ and for each e , $a_{e,t} = a$ for some different a .

Proof

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

Let $v(\lambda, \gamma) := (u(\lambda, \gamma) - u(e, \gamma))_{e \in E} \in \mathbb{R}^{|E|}$ and $\bar{v}_T := \frac{1}{T} \sum_{t \leq T} v(\sigma(h_t), \sigma_0(h_t))$. 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 $x \in C$.

Application: Picking Experts – Proof (via Blackwell)

Proof

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

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

Application: Picking Experts – Proof (via Blackwell)

Proof

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

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

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

$$(i) \neg(x \geq 0) \implies \|x^-\| = \|x - y\| > 0 \implies x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y);$$

Application: Picking Experts – Proof (via Blackwell)

Proof

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

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

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

(i) $\neg(x \geq 0) \implies \|x^-\| = \|x - y\| > 0 \implies x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y)$;

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

Note that

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

Application: Picking Experts – Proof (via Blackwell)

Proof

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

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

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

(i) $\neg(x \geq 0) \implies \|x^-\| = \|x - y\| > 0 \implies x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y)$;

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

Note that

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

Hence C is a B -set. Blackwell's theorem $\implies \bar{v}_t$ approaches C a.s., i.e.,

$\liminf_{t \rightarrow \infty} (\bar{u}_t(\sigma, \sigma_0) - \bar{u}_t(e, \sigma_0)) \geq 0$ for all $e \in E$ and any of nature's moves σ_0 .

Application: Picking Experts – Proof (via Blackwell)

Proof

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

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

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

(i) $\neg(x \geq 0) \implies \|x^-\| = \|x - y\| > 0 \implies x \in H^+(x - y, (x - y) \cdot y) \setminus H(x - y, (x - y) \cdot y)$;

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

Note that

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

Hence C is a B -set. Blackwell's theorem $\implies \bar{v}_t$ approaches C a.s., i.e.,

$\liminf_{t \rightarrow \infty} (\bar{u}_t(\sigma, \sigma_0) - \bar{u}_t(e, \sigma_0)) \geq 0$ for all $e \in E$ and any of nature's moves σ_0 .

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

Overview

1. Learning in Games

2. Approachability

- Multi-Utility Representation
- Approachability
- Application: Picking Experts
- **Universal Consistency**

3. Calibration

4. Adaptive Algorithms

5. Sophisticated Learning

Hannan (1957): Setting

Repeated decision problem

Actions $A = \{1, \dots, |A|\}$; states $S = \{1, \dots, |S|\}$.

Bounded stage payoff $u : A \times S \rightarrow [0, 1]$.

At t : DM chooses $a_t \in A$; Nature reveals $s_t \in S$; payoff $u(a_t, s_t)$.

Hannan (1957): Setting

Repeated decision problem

Actions $A = \{1, \dots, |A|\}$; states $S = \{1, \dots, |S|\}$.

Bounded stage payoff $u : A \times S \rightarrow [0, 1]$.

At t : DM chooses $a_t \in A$; Nature reveals $s_t \in S$; payoff $u(a_t, s_t)$.

Beliefs and empirical distribution

$\hat{p}_t \in \Delta(S)$: empirical distribution of $(s_\ell)_{\ell \leq t}$.

Bayes payoff of action a against $p \in \Delta(S)$: $U(a, p) := \mathbb{E}_{s \sim p}[u(a, s)]$.

Repeated decision problem

Actions $A = \{1, \dots, |A|\}$; states $S = \{1, \dots, |S|\}$.

Bounded stage payoff $u : A \times S \rightarrow [0, 1]$.

At t : DM chooses $a_t \in A$; Nature reveals $s_t \in S$; payoff $u(a_t, s_t)$.

Beliefs and empirical distribution

$\hat{p}_t \in \Delta(S)$: empirical distribution of $(s_\ell)_{\ell \leq t}$.

Bayes payoff of action a against $p \in \Delta(S)$: $U(a, p) := \mathbb{E}_{s \sim p}[u(a, s)]$.

Benchmark: (External) Regret

For any fixed $a \in A$, $\bar{u}_T(a) := \frac{1}{T} \sum_{t \leq T} u(a, s_t)$.

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

No (External) Regret: $\liminf_{T \rightarrow \infty} (\bar{u}_T(\sigma) - \max_a \bar{u}_T(a)) \geq 0$ a.s.

External regret: comparison relative to swapping to fixed action.

Smoothed fictitious play

Fix a full-support $\mathbf{v}_t \in \Delta(S)$ iid and a sequence (γ_t) with $\gamma_t \downarrow 0$ and $\sum_t \gamma_t < \infty$.

At $t \geq 1$, form the **smoothed empirical belief**

$$p_t := (1 - \gamma_t) \hat{p}_{t-1} + \gamma_t \mathbf{v}_t.$$

Choose a Bayes action (or mixed Bayes rule) against p_t :

$$a_t \in \arg \max_{a \in A} u(a, p_t) \quad \left(\text{or } \lambda_t \in \arg \max_{\lambda \in \Delta(A)} \sum_a \lambda(a) u(a, p_t) \right).$$

Ties are broken by a fixed rule; γ_t prevents zero exploration.

Interpretation: best respond to slightly perturbed empirical model;
perturbations vanish.

Theorem (Hannan 1957)

Under SFP with bounded payoffs and (γ_t) as above,

$$\liminf_{T \rightarrow \infty} \left(\bar{u}_T(\sigma) - \max_{a \in A} \bar{u}_T(a) \right) \geq 0 \quad \text{almost surely.}$$

Equivalently, for every $a \in A$, $\liminf_{T \rightarrow \infty} (\bar{u}_T(\sigma) - \bar{u}_T(a)) \geq 0$ a.s.

Full-support $\mathbf{v}_t$ ensures absolute continuity of beliefs; (γ_t) controls approximation error.

Key steps

Let $v_t(a) := u(a_t, s_t) - u(a, s_t)$; regret against a is $\bar{v}_T(a) = \bar{u}_T(\sigma) - \bar{u}_T(a)$.

Conditional on H_{t-1} , a_t maximises $u(\cdot, p_t)$, hence for all a ,

$$\mathbb{E}[v_t(a) \mid H_{t-1}] = u(a_t, p_t) - u(a, p_t) + \underbrace{\left(\mathbb{E}[u(a_t, s_t) - u(a, s_t) \mid H_{t-1}] - [u(a_t, p_t) - u(a, p_t)] \right)}_{\text{model error}}.$$

The first term ≥ 0 by optimality of a_t for p_t .

The model-error term is $O(\gamma_t)$ since $p_t - (1 - \gamma_t)\hat{p}_{t-1}$ has mass γ_t and payoffs are bounded in $[0, 1]$.

Summing and dividing by T gives $\mathbb{E}[\bar{v}_T(a)] \geq -\frac{1}{T} \sum_{t \leq T} c \gamma_t$; with $\sum_t \gamma_t < \infty$,
 $\liminf_T \mathbb{E}[\bar{v}_T(a)] \geq 0$.

A standard martingale SLLN upgrades to a.s. statements (bounded differences).

Direct No-Regret via Approachability

History $H = \cup_t H_t$, $H_t := (A \times S)^{t-1}$. Strategy $\sigma : H \rightarrow \Delta(A)$; Nature $\sigma_0 : H \rightarrow \Delta(S)$.

Define **regret vector** $v(\lambda, \gamma) \in \mathbb{R}^A$ with $v(\lambda, \gamma) := (u(\lambda, \gamma) - u(a, \gamma))_{a \in A}$.

$$\bar{v}_T := \frac{1}{T} \sum_{t \leq T} v(\sigma(h_t), \sigma_0(h_t)).$$

Target set $C = \mathbb{R}_+^A$ (no external regret).

Let $x \in \mathbb{R}^{|A|}$, $y := x^+$, $x^- := y - x$, and choose

$$\lambda^x(a) = \frac{x_a^-}{\sum_{a'} x_{a'}^-} \quad \text{if } x \notin C, \quad \text{any } \lambda^x \text{ if } x \in C.$$

Then $x^- \cdot \mathbb{E}[r_t \mid H_{t-1}] = \mathbf{0} \leq \mathbf{0}$; C is a B -set $\implies \bar{v}_T \rightarrow C$ a.s.

Consequence: $\liminf_T (\bar{u}_T(\sigma) - \bar{u}_T(a)) \geq \mathbf{0}$ for all a .

Direct No-Regret via Approachability

History $H = \cup_t H_t$, $H_t := (A \times S)^{t-1}$. Strategy $\sigma : H \rightarrow \Delta(A)$; Nature $\sigma_0 : H \rightarrow \Delta(S)$.

Define **regret vector** $v(\lambda, \gamma) \in \mathbb{R}^A$ with $v(\lambda, \gamma) := (u(\lambda, \gamma) - u(a, \gamma))_{a \in A}$.

$$\bar{v}_T := \frac{1}{T} \sum_{t \leq T} v(\sigma(h_t), \sigma_0(h_t)).$$

Target set $C = \mathbb{R}_+^A$ (no external regret).

Let $x \in \mathbb{R}^{|A|}$, $y := x^+$, $x^- := y - x$, and choose

$$\lambda^x(a) = \frac{x_a^-}{\sum_{a'} x_{a'}^-} \quad \text{if } x \notin C, \quad \text{any } \lambda^x \text{ if } x \in C.$$

Then $x^- \cdot \mathbb{E}[r_t \mid H_{t-1}] = \mathbf{0} \leq \mathbf{0}$; C is a B -set $\implies \bar{v}_T \rightarrow C$ a.s.

Consequence: $\liminf_T (\bar{u}_T(\sigma) - \bar{u}_T(a)) \geq \mathbf{0}$ for all a .

Remark

Hannan (1957) and Blackwell (1956) yield same external no-regret guarantee; update rules differ (smoothed best reply vs projection-based mixture).

Choice of Smoothing and Rates:

Any full-support $\mathbf{v}_t$ works; e.g. $\mathbf{v}$ uniform on S .

Typical schedules: $\gamma_t = t^{-1-\epsilon}$ (summable) for a.s. convergence via almost-supermartingale; or $\gamma_t \asymp t^{-1/2}$ for $O(1/\sqrt{T})$ expected regret.

Boundedness of u and tie-breaking rules ensure measurability and martingale applicability.

Implementation, Interpretation and Links

Choice of Smoothing and Rates:

Any full-support $\mathbf{v}_t$ works; e.g. $\mathbf{v}$ uniform on S .

Typical schedules: $\gamma_t = t^{-1-\epsilon}$ (summable) for a.s. convergence via almost-supermartingale; or $\gamma_t \asymp t^{-1/2}$ for $O(1/\sqrt{T})$ expected regret.

Boundedness of u and tie-breaking rules ensure measurability and martingale applicability.

Exploration: $\gamma_t > 0$ avoids being trapped by early noise; $\gamma_t \downarrow 0$ removes bias.

Interpretation: “Fictitious play” against the environment with vanishing perturbations.

Implementation, Interpretation and Links

Choice of Smoothing and Rates:

Any full-support $\mathbf{v}_t$ works; e.g. $\mathbf{v}$ uniform on S .

Typical schedules: $\gamma_t = t^{-1-\epsilon}$ (summable) for a.s. convergence via almost-supermartingale; or $\gamma_t \asymp t^{-1/2}$ for $O(1/\sqrt{T})$ expected regret.

Boundedness of u and tie-breaking rules ensure measurability and martingale applicability.

Exploration: $\gamma_t > 0$ avoids being trapped by early noise; $\gamma_t \downarrow 0$ removes bias.

Interpretation: “Fictitious play” against the environment with vanishing perturbations.

Connections: Approachability (Blackwell), SFP.

Scope: applies to arbitrary state sequences (adversarial or stochastic).

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Foster and Vohra (1997): Calibrated Learning & CE

4. Adaptive Algorithms

5. Sophisticated Learning

Overview

1. Learning in Games

2. Approachability

3. Calibration

- Foster and Vohra (1997): Calibrated Learning & CE

4. Adaptive Algorithms

5. Sophisticated Learning

Learning in Games

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.

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.

Calibrated Learning: Setup

Stage game

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

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

Repeated play ($t = 1, 2, \dots$)

History $H_t := A^{t-1}$, $H := \cup_t H_t$.

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

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

Forecasts and behaviour

Player i 's forecasting rule $f_{-i}^i : H \rightarrow \Delta(A_{-i})$; issues forecasts $f_{-i}^i(h_t) = \sigma_{-i,t}^i \in \Delta(A_{-i})$.

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

Fix a deterministic tie-breaking rule.

Calibration (Partition Version)

Definition

Fix finite partition $\Pi_i = \{B_{-i}^k\}_{k=1}^K$ of $\Delta(A_{-i})$ and representative $\sigma_{-i}^k \in B_{-i}^k$. For $t \geq 1$ set

$$N_{i,t}^k := \sum_{s \leq t} \mathbf{1}_{\{\sigma_{-i,s}^i \in B_{-i}^k\}}, \quad \bar{\sigma}_{-i,t}^k(a_{-i}) := \frac{1}{N_{i,t}^k} \sum_{s \leq t} \mathbf{1}_{\{\sigma_{-i,s}^i \in B_{-i}^k\}} \mathbf{1}_{\{a_{-i,s} = a_{-i}\}} \text{ if } N_{i,t}^k > 0; \text{ ow } = 0.$$

The forecasting rule is **calibrated (wrt $(a_{-i,t})$ on Π_i)** if for every k

$$\lim_{t \rightarrow \infty} \|\bar{\sigma}_{-i,t}^k - \sigma_{-i}^k\| \frac{N_{i,t}^k}{t} = 0.$$

Intuition: on subsequence when $\sigma_{-i,t}^i \in B_{-i}^k$, empirical conditional frequency of $a_{-i,t}$ matches σ_{-i}^k .

Refining Π_i (mesh $\downarrow 0$) yields full calibration on $\Delta(A_{-i})$.

Existence of Calibrated Forecasters

Proposition (Foster and Vohra 1997 GEB)

For any finite partition Π_i there exists a (possibly randomised) forecasting scheme $f_{-i}^i : H \rightarrow \Delta(A_{-i})$ that is calibrated on Π_i a.s.

Proof idea (via Blackwell approachability)

Build *calibration vector* $z_t \in \mathbb{R}^{K|A_{-i}|}$ with components

$$z_t(k, a_{-i}) := \mathbf{1}_{\{\sigma_{-i,t}^j \in B_{-i}^k\}} (\mathbf{1}_{\{a_{-i,t}=a_{-i}\}} - \sigma_{-i}^k(a_{-i})).$$

Calibration vector average: $\bar{z}_t := \sum_{s \leq t} z_s / t$

Target set $C := \mathbb{R}_{-}^{K|A_{-i}|}$.

Note that $d(\bar{z}_t, C) \rightarrow 0 \iff \bar{z}_\infty(k, a_{-i}) \leq \sigma_{-i}^k(a_{-i}) \forall a_{-i} \text{ and } k : \lim_{t \rightarrow \infty} N_{i,t}^k / t > 0$
 $\iff \bar{z}_\infty(k, a_{-i}) = \sigma_{-i}^k(a_{-i})$.

At step t , choose a forecast cell k (i.e. $\sigma_{-i,t}^j \in B_{-i}^k$) to satisfy Blackwell's separation for the current average $\bar{z}_{t-1}$.

Blackwell $\implies \bar{z}_t \rightarrow 0$ a.s. on each active cell $k \implies$ calibration on Π_i .

Calibration via Blackwell Approachability: Vector Game

Fix finite forecast menu $S = \{\sigma_{-i}^1, \dots, \sigma_{-i}^K\} \subset \Delta(A_{-i})$ that *contains all pure actions* $\{e_{a_{-i}}\}$.

At t , forecaster chooses index $k_t \in \{1, \dots, K\}$ (announces $\sigma_{-i}^{k_t}$); Nature chooses $a_{-i,t} \in A_{-i}$.

Stage vector payoff $z_t \in \mathbb{R}^{K|A_{-i}|}$ with blocks $z_t(k, \cdot) \in \mathbb{R}^{|A_{-i}|}$:

$$z_t(k, a_{-i}) := \mathbf{1}_{\{k_t=k\}} \left(\mathbf{1}_{\{a_{-i,t}=a_{-i}\}} - \sigma_{-i}^k(a_{-i}) \right).$$

Averaging gives (for each k)

$$\bar{z}_t(k, \cdot) = \frac{N_t(k)}{t} (\bar{\sigma}_t^k - \sigma_{-i}^k), \quad N_t(k) := \sum_{s \leq t} \mathbf{1}_{\{k_s=k\}}, \quad \bar{\sigma}_t^k(a_{-i}) := \frac{\sum_{s \leq t} \mathbf{1}_{\{k_s=k\}} \mathbf{1}_{\{a_{-i,s}=a_{-i}\}}}{N_t(k)}.$$

Target set $C := \{0\} \subset \mathbb{R}^{K|A_{-i}|}$.

Goal: $\bar{z}_t \rightarrow 0$ a.s. $\iff$ for each k , $\|\bar{\sigma}_t^k - \sigma_{-i}^k\| \cdot N_t(k)/t \rightarrow 0$ (partition calibration on S).

Separation and Control Rule (choose k_t)

Separation at current average

Let $x := \bar{z}_{t-1} \in \mathbb{R}^{K|A_{-i}|}$ and write its k -block as $x_k \in \mathbb{R}^{|A_{-i}|}$.

If we announce cell k and Nature uses $\lambda_{-i} \in \Delta(A_{-i})$, then

$$\mathbb{E}[z_t \mid H_{t-1}, k, \lambda_{-i}] \cdot x = (\lambda_{-i} - \sigma_{-i}^k) \cdot x_k.$$

Choose an index k_t such that *its fixed representative is a maximiser of the linear form $q \mapsto x_{k_t} \cdot q$ over the whole simplex*:

$$\sigma_{-i}^{k_t} \in \arg \max_{q \in \Delta(A_{-i})} x_{k_t} \cdot q.$$

This exists because a linear form on $\Delta(A_{-i})$ is maximised at an extreme point; since all pure actions are in S , we can take $\sigma_{-i}^{k_t} = e_{j^*}$ where $j^* \in \arg \max_{a_{-i}} x_{k_t}(a_{-i})$.

Then for every λ_{-i} ,

$$\sup_{\lambda_{-i} \in \Delta(A_{-i})} \mathbb{E}[z_t \mid H_{t-1}, k_t, \lambda_{-i}] \cdot x = \left(\max_{\lambda_{-i}} x_{k_t} \cdot \lambda_{-i} \right) - x_{k_t} \cdot \sigma_{-i}^{k_t} \leq 0.$$

Conclusion

The origin $C = \{0\}$ satisfies Blackwell's separation condition for this control rule. Hence

Since C and the forecaster are both approachable, the forecaster can achieve a vanishing regret.

From Approachability to Calibration (on S)

Blackwell's theorem (bounded vectors) $\implies \bar{z}_t \rightarrow \mathbf{0}$ a.s.

For each k :

$$\bar{z}_t(k, \cdot) = \frac{N_t(k)}{t} (\hat{\lambda}_t^k - \sigma_{-i}^k) \xrightarrow[t \rightarrow \infty]{} \mathbf{0},$$

hence $\|\hat{\lambda}_t^k - \sigma_{-i}^k\| \cdot \frac{N_t(k)}{t} \rightarrow \mathbf{0}$.

Calibration (partition version): the forecast is calibrated w.r.t. the (finite) forecast menu S actually used.

If both players best reply to their calibrated forecasts, standard Foster–Vohra argument $\implies$ the empirical distribution of play is a (coarse) correlated equilibrium; with vanishing mesh grids, obtain correlated equilibrium.

Implementation Recipe (what to compute each period)

At the start of period t

Maintain $x = \bar{z}_{t-1}$ and its blocks x_k .

For each k whose representative is pure e_a , compute $x_k(a)$.

Pick k_t with representative $\sigma_{-j}^{k_t} = e_{a^*}$ where $a^* \in \arg \max_a x_{k_t}(a)$.

Announce forecast $\sigma_{-i,t}^j := \sigma_{-j}^{k_t}$; play a myopic best reply to this forecast (for the “learning $\rightarrow$ CE” part).

Update z_t and the running average.

Remark

To move from calibration on a finite menu S to ϵ -calibration on $\Delta(A_{-j})$, take S to be a fine grid (mesh $\leq \epsilon$) *together with all pure actions*. The same separation rule works.

Existence of Calibrated Forecasters

Separation step and control rule

Let $x := \bar{z}_{t-1} \in \mathbb{R}^{K|A_{-i}|}$. Target set $C := \mathbb{R}_{-}^{K|A_{-i}|}$. Projection $y := \mathbf{0} \wedge x$; normal $x - y = x^+$.

If we announce cell k (i.e. forecast σ_{-i}^k) and Nature uses $\lambda_{-i} \in \Delta(A_{-i})$, then

$$\mathbb{E}[z_t \mid H_{t-1}, k, \lambda_{-i}] = e_k \otimes (\lambda_{-i} - \sigma_{-i}^k),$$

where e_k is k -th basis vector and $\otimes$ concatenates $|A_{-i}|$ -block at k .

For Blackwell's approachability, need: $\mathbb{E}[z_t \mid H_{t-1}, k, \lambda_{-i}] \cdot x^+ = (\lambda_{-i} - \sigma_{-i}^k) \cdot x_k^+ \leq 0$,

$\forall \lambda_{-i} \in \Delta(A_{-i})$, where x_k^+ is the k -block of x^+ . This holds iff $\sigma_{-i}^k \in \arg \max_{q \in \Delta(A_{-i})} x_k^+ \cdot q$.

Rule at time t : compute x_k^+ for each k and choose

$$k_t \in \arg \max_k x_k^+ \cdot \sigma_{-i}^k, \quad \text{then forecast } \sigma_{-i,t}^i := \sigma_{-i}^{k_t}.$$

Then for every λ_{-i} , $(\lambda_{-i} - \sigma_{-i}^{k_t}) \cdot x_k^+ \leq 0 \implies$ Blackwell separation at x .

Implementation remark

Include all pure actions among representatives $\{\sigma_{-i}^k\}$; since $\lambda_{-i} \mapsto x_k^+ \cdot \lambda_{-i}$ is linear, maximiser can be taken pure, ensuring argmax is available.

Correlated Equilibrium

Definition (correlated equilibrium)

$\pi \in \Delta(A)$ is a **correlated equilibrium** if for all maps $F_1 : A_1 \rightarrow A_1$ and $F_2 : A_2 \rightarrow A_2$,

$$\sum_a \pi(a) [u_1(a) - u_1(F_1(a_1), a_2)] \geq 0, \quad \sum_a \pi(a) [u_2(a) - u_2(a_1, F_2(a_2))] \geq 0.$$

Equivalent to Aumann's "no profitable deviation conditional on signal".

Denote the set by CE .

Theorem (Foster and Vohra 1997)

Suppose each player uses a calibrated forecasting scheme (on arbitrarily fine partitions) and in each t plays a myopic best reply to their forecast (fixed tie-breaking). Then every limit point of (D_t) lies in $\mathcal{CE}$.

Proof

For player 1, define best-reply regions $M_1(a_1) := \{q \in \Delta(A_2) : a_1 \in \arg \max_{x_1} \sum_{a_2} q(a_2) u_1(x_1, a_2)\}$.

By tie-breaking, whenever $p_t \in B^k \subseteq M_1(a_1)$, player 1 plays a_1 .

Calibration on $\Pi \implies$ conditional empirical law of $a_{2,t}$ given $\{p_t \in B^k\}$ converges to q^k ;
hence any limit π satisfies $\pi(\cdot \mid a_1) \in M_1(a_1)$ whenever $\pi(a_1) > 0$.

Symmetrically for player 2: $\pi(\cdot \mid a_2) \in M_2(a_2)$ if $\pi(a_2) > 0$.

These conditions are equivalent to the CE inequalities $\implies \pi \in \mathcal{CE}$.

Attainability of Any Correlated Equilibrium

Theorem (Foster and Vohra 1997)

For any $\pi^* \in CE$ there exist calibrated forecasters and myopic best replies such that $D_t \rightarrow \pi^*$ (all limit points equal π^*).

Construct partitions and representatives matching π^* 's conditionals; calibrate to them.

Best replies implement the recommended supports; empirical play tracks π^* .

From Calibration to CE: Details

Write $D_t(a_1, \cdot)$ for row- a_1 . If $\liminf_t \sum_{a_2} D_t(a_1, a_2) > 0$, then along any convergent subsequence $D_{t_k} \rightarrow \pi$,

$$\frac{D_{t_k}(a_1, \cdot)}{\sum_{a_2} D_{t_k}(a_1, a_2)} \rightarrow q^{a_1} \in M_1(a_1).$$

Hence $\sum_{a_2} \pi(a_1, a_2) [u_1(a_1, a_2) - u_1(x_1, a_2)] \geq 0$ for all x_1 .

Do the same for player 2; collect the inequalities to obtain the CE conditions.

Constructing Calibrated Forecasts (Procedure)

Fix partition $\Pi = \{B^k\}$ with representatives (q^k) .

Maintain running average $\bar{z}_{t-1} \in \mathbb{R}^{K|A_2|}$ of calibration vectors $z_t(k, a_2)$.

Forecast rule (Blackwell step):

If $\bar{z}_{t-1} \in C$ (all active components near $\mathbf{0}$), pick any k .

Else let $x^- := (\bar{z}_{t-1})^-$; choose k that minimises $\sum_{a_2} x^-(k, a_2) q^k(a_2)$; set $p_t \in B^k$.

Guarantees: $\bar{z}_t \rightarrow \mathbf{0}$ a.s. on all active $k \implies$ calibration on Π .

Refining Π over time (mesh $\downarrow \mathbf{0}$) yields full calibration.

Interpretation and Links

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 (contrast with calibration-based route).

Design/selection: by designing calibrated grids, any target $\pi^* \in CE$ can be attained.

Uniform Calibration (Refinement Limit)

Definition (uniform ε -calibration)

A forecast sequence (p_t) is **ε -calibrated** if there exist points $q^1, \dots, q^K$ with $\max_{a_2} \min_k |p(a_2) - q^k(a_2)| \leq \varepsilon$ for all $p \in \Delta(A_2)$ and

$$\limsup_{t \rightarrow \infty} \sum_{k=1}^K \frac{N^k(t)}{t} \max_{a_2} |r_t^k(a_2) - q^k(a_2)| \leq \varepsilon.$$

Existence for all $\varepsilon > 0$; letting $\varepsilon \downarrow 0$ yields full calibration.

Compatible with the approachability construction by refining the partition.

Calibrated learning procedure: learning procedure such that in the long run each action is a best response to the frequency distribution of opponents' choices in all periods in which that action was played

Foster Vohra 1997 GEB, Calibrated Learning and Correlated Equilibrium

Foster Hart 2018 GEB, Smooth calibration, leaky forecasts, finite recall, and Nash dynamics

Foster Hart 2021 JPE, Forecast Hedging and Calibration

The first study of no-regret strategies was conducted by Hannan [1957]. The connection between no-regret strategies and the concept of approachable sets was first made by Hart and Mas-Colell [2000]. Several studies, including Foster and Vohra [1997] and Fudenberg and Levine [1999], define no-regret in a stronger form than the one presented here. Rustichini [1999], Lugosi, Mannor, and Stoltz [2007], and Lehrer and Solan [2007] studied no-regret strategies under which the decision maker does not know the true state of nature, but receives information that depends on the state of nature and the chosen action.

Overview

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

Papers: Hart Mas-Colell 2003 AER, Uncoupled Dynamics Do Not Lead to Nash Equilibrium

*Hart 2005 Ect, Adaptive Heuristics

Foster Young 2006 TE, Regret testing. learning to play Nash equilibrium without knowing you have an opponent

Papers on reinforcement learning and Q-learning

What is players are Bayesian wrt gameplay and engage in sophisticated learning?

Sophisticated Learning

What is players are Bayesian wrt gameplay and engage in sophisticated learning?

Two papers:

Kalai and Lehrer (1993 Ecta) "Rational Learning Leads to Nash Equilibria"

Kalai and Lehrer (1993 Ecta) "Subjective Equilibrium in Repeated Games"

(Will favour Fudenberg and Levine's "sophisticated learning" terminology.)

Stage Game and Repeated Interaction

Players $i \in I = \{1, \dots, n\}$; actions A_i (finite). Profile $A = \times_i A_i$.

Payoffs $u_i : A \rightarrow \mathbb{R}$. One-period outcome $a^t = (a_i^t)_i \in A$.

Stage Game and Repeated Interaction

Players $i \in I = \{1, \dots, n\}$; actions A_i (finite). Profile $A = \times_i A_i$.

Payoffs $u_i : A \rightarrow \mathbb{R}$. One-period outcome $a^t = (a_i^t)_i \in A$.

Repeated game: infinite horizon, perfect monitoring, discounts $\delta_i \in (0, 1)$.

Histories $h^t = (a^0, \dots, a^{t-1}) \in H^t := A^t$; $H = \cup_{t \geq 0} H^t$; $\emptyset$ at $t = 0$.

Behavioural strategies $\sigma_i = (\sigma_{i,t})_{t \geq 0}$, with $\sigma_{i,t} : H^t \rightarrow \Delta(A_i)$.

Strategy profile $\sigma = (\sigma_i)_i$. Outcome law μ^σ on $\Omega := A^{\mathbb{N}}$ (product σ -algebra).

Stage Game and Repeated Interaction

Players $i \in I = \{1, \dots, n\}$; actions A_i (finite). Profile $A = \times_i A_i$.

Payoffs $u_i : A \rightarrow \mathbb{R}$. One-period outcome $a^t = (a_i^t)_i \in A$.

Repeated game: infinite horizon, perfect monitoring, discounts $\delta_i \in (0, 1)$.

Histories $h^t = (a^0, \dots, a^{t-1}) \in H^t := A^t$; $H = \cup_{t \geq 0} H^t$; $\emptyset$ at $t = 0$.

Behavioural strategies $\sigma_i = (\sigma_{i,t})_{t \geq 0}$, with $\sigma_{i,t} : H^t \rightarrow \Delta(A_i)$.

Strategy profile $\sigma = (\sigma_i)_i$. Outcome law μ^σ on $\Omega := A^\mathbb{N}$ (product σ -algebra).

History concatenation: $hh' \in H^{t+r} : h \in H^t, h' \in H^r$.

Continuation histories starting from h_t : $C(h_t) := \{h' \in H^\infty \mid (h_t h') \in H^\infty\}$.

Filtration $(\mathcal{F}_t)$, $\mathcal{F}_t := \sigma(\{h^t\})$.

Normalised expected discounted payoff:

$$U_i(\sigma) = (1 - \delta_i) \mathbb{E}_{\mu^\sigma} \left[\sum_{t \geq 0} \delta_i^t u_i(a^t) \right].$$

Beliefs, Absolute Continuity, and Payoffs

Player i 's conjectures/degenerate beliefs about opponents' strategies σ_{-i}^i .

Induces belief $\mu_i = \mu^{\sigma_{-i}^i}$ on Ω .

Player i 's prior ν_i on opponents' strategies σ_{-i} (Actual uncertainty).

Induces belief μ_i on Ω via $\tilde{\sigma}_{-i} \mapsto \mu^{(\sigma_i, \tilde{\sigma}_{-i})}$.

For ν_i , expected conjecture: $\sigma_{-i}^i(h)(a_{-i}) = \mathbb{E}_{\tilde{\sigma}_{-i} \sim \nu_i}[\tilde{\sigma}_{-i}(h)(a_{-i})]$.

Player i 's **Subjective joint strategy**: $\sigma^i = (\sigma_i, \sigma_{-i}^i)$.

Beliefs, Absolute Continuity, and Payoffs

Player i 's conjectures/degenerate beliefs about opponents' strategies σ_{-i}^i .

Induces belief $\mu_i = \mu^{\sigma_{-i}^i}$ on Ω .

Player i 's prior ν_i on opponents' strategies σ_{-i} (Actual uncertainty).

Induces belief μ_i on Ω via $\tilde{\sigma}_{-i} \mapsto \mu^{(\sigma_i, \tilde{\sigma}_{-i})}$.

For ν_i , expected conjecture: $\sigma_{-i}^i(h)(a_{-i}) = \mathbb{E}_{\tilde{\sigma}_{-i} \sim \nu_i}[\tilde{\sigma}_{-i}(h)(a_{-i})]$.

Player i 's **Subjective joint strategy**: $\sigma^i = (\sigma_i, \sigma_{-i}^i)$.

Truth-compatibility (absolute continuity): $\mu^\sigma \ll \mu_i$ for all i .

(i.e., $\mu^\sigma(E) > 0 \implies \mu_i(E) > 0$ for any μ_i -measurable E .)

Posteriors: after h^t , update $\mu_i(\cdot \mid h^t)$ by Bayes (well-defined by abs. cont.).

Beliefs, Absolute Continuity, and Payoffs

Player i 's conjectures/degenerate beliefs about opponents' strategies σ_{-i}^i .

Induces belief $\mu_i = \mu^{\sigma_{-i}^i}$ on Ω .

Player i 's prior ν_i on opponents' strategies σ_{-i} (Actual uncertainty).

Induces belief μ_i on Ω via $\tilde{\sigma}_{-i} \mapsto \mu^{(\sigma_i, \tilde{\sigma}_{-i})}$.

For ν_i , expected conjecture: $\sigma_{-i}^i(h)(a_{-i}) = \mathbb{E}_{\tilde{\sigma}_{-i} \sim \nu_i}[\tilde{\sigma}_{-i}(h)(a_{-i})]$.

Player i 's **Subjective joint strategy**: $\sigma^i = (\sigma_i, \sigma_{-i}^i)$.

Truth-compatibility (absolute continuity): $\mu^\sigma \ll \mu_i$ for all i .

(i.e., $\mu^\sigma(E) > 0 \implies \mu_i(E) > 0$ for any μ_i -measurable E .)

Posteriors: after h^t , update $\mu_i(\cdot \mid h^t)$ by Bayes (well-defined by abs. cont.).

Rationality path: each period t , $\sigma_{i,t}$ is a best response to $\mu_i(\cdot \mid h^t)$.

Induced strategy: for histories $h, h' \in H$, denote $\sigma_h(h') := \sigma(hh')$ (strategy following h for h').

Closeness and “Plays ϵ -Like”

Definition (ϵ -close measures)

For $\epsilon > 0$, μ is ϵ -**close** to $\tilde{\mu}$ if $\exists Q$ with $\mu(Q), \tilde{\mu}(Q) \geq 1 - \epsilon$ s.t. $\forall$ measurable $A \subseteq Q$,

$$(1 - \epsilon)\tilde{\mu}(A) \leq \mu(A) \leq (1 + \epsilon)\tilde{\mu}(A).$$

Closeness and “Plays ϵ -Like”

Definition (ϵ -close measures)

For $\epsilon > 0$, μ is **ϵ -close** to $\tilde{\mu}$ if $\exists Q$ with $\mu(Q), \tilde{\mu}(Q) \geq 1 - \epsilon$ s.t. $\forall$ measurable $A \subseteq Q$,

$$(1 - \epsilon)\tilde{\mu}(A) \leq \mu(A) \leq (1 + \epsilon)\tilde{\mu}(A).$$

Definition (plays ϵ -like)

A profile σ **plays ϵ -like** σ' if μ^σ is ϵ -close to $\mu^{\sigma'}$; equivalently, after any h^t , the conditional laws are ϵ -close on a large-probability subset.

Closeness and “Plays ϵ -Like”

Definition (ϵ -close measures)

For $\epsilon > 0$, μ is ϵ -**close** to $\tilde{\mu}$ if $\exists Q$ with $\mu(Q), \tilde{\mu}(Q) \geq 1 - \epsilon$ s.t. $\forall$ measurable $A \subseteq Q$,

$$(1 - \epsilon)\tilde{\mu}(A) \leq \mu(A) \leq (1 + \epsilon)\tilde{\mu}(A).$$

Definition (plays ϵ -like)

A profile σ **plays ϵ -like** σ' if μ^σ is ϵ -close to $\mu^{\sigma'}$; equivalently, after any h^t , the conditional laws are ϵ -close on a large-probability subset.

Controls conditional probabilities on tails; prevents cumulative small-error blowup across time.

Theorem 1 (Learning to predict)

Fix actual strategy σ and player i 's subjective joint strategy $\sigma^i := (\sigma_i, \sigma_{-i}^i)$. If $\mu^\sigma \ll \mu^{\sigma^i}$, then for every $\varepsilon > 0$ and for μ^σ -a.e. path $h \in H^\infty$, $\exists T$ s.t. $\forall t \geq T$, continuation σ_{h_t} plays ε -like $\sigma_{h_t}^i$.

Theorem 1 (Learning to predict)

Fix actual strategy σ and player i 's subjective joint strategy $\sigma^i := (\sigma_i, \sigma_{-i}^i)$. If $\mu^\sigma \ll \mu^{\sigma^i}$, then for every $\varepsilon > 0$ and for μ^σ -a.e. path $h \in H^\infty$, $\exists T$ s.t. $\forall t \geq T$, continuation σ_{h_t} plays ε -like $\sigma_{h_t}^i$.

Posterior forecasts of future play (conditional on realised history) merge with truth.

No optimality required here; this is a property of Bayesian updating under abs. cont.

Theorem 3 (Blackwell and Dubins, 1962)

If $\mu \ll \tilde{\mu}$, then with μ -probability 1, for every $\varepsilon > 0$ there exists random time $\tau(\varepsilon)$ such that for all $t \geq \tau(\varepsilon)$ the posteriors $\mu(\cdot \mid \mathcal{F}_t)$ and $\tilde{\mu}(\cdot \mid \mathcal{F}_t)$ are ε -close.

Theorem 3 (Blackwell and Dubins, 1962)

If $\mu \ll \tilde{\mu}$, then with μ -probability 1, for every $\varepsilon > 0$ there exists random time $\tau(\varepsilon)$ such that for all $t \geq \tau(\varepsilon)$ the posteriors $\mu(\cdot \mid \mathcal{F}_t)$ and $\tilde{\mu}(\cdot \mid \mathcal{F}_t)$ are ε -close.

If people start off with compatible priors, posteriors become arbitrarily close after exposed to enough information.

Merging via Likelihood Ratios

Theorem 3 (Blackwell and Dubins, 1962)

If $\mu \ll \tilde{\mu}$, then with μ -probability 1, for every $\varepsilon > 0$ there exists random time $\tau(\varepsilon)$ such that for all $t \geq \tau(\varepsilon)$ the posteriors $\mu(\cdot \mid \mathcal{F}_t)$ and $\tilde{\mu}(\cdot \mid \mathcal{F}_t)$ are ε -close.

If people start off with compatible priors, posteriors become arbitrarily close after exposed to enough information.

Proof Idea

Radon-Nikodym derivative $\phi = \frac{d\mu}{d\tilde{\mu}}$ exists; set $M_t = \mathbb{E}_{\tilde{\mu}}[\phi \mid \mathcal{F}_t]$.

(M_t) is a nonnegative $\tilde{\mu}$ -martingale; $M_t \rightarrow M_\infty$ a.s.

Control likelihood ratios on Q with $\mu(Q), \tilde{\mu}(Q) \approx 1$.

Translate bounds to conditionals on continuation histories $C(h^t)$; conclude ε -closeness.

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

σ_i is a best response to σ_{-i}^i , for every i ;

σ plays ε -like σ^i , for every i .

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

σ_i is a best response to σ_{-i}^i , for every i ;

σ plays ε -like σ^i , for every i .

Corollary 1

If each σ_i best responds to σ_{-i}^i and $\sigma \ll \sigma^i$ for all i , then for a.e. path $h \exists T$ s.t. $\forall t \geq T$, the continuation σ_{h_t} is a subjective ε -equilibrium.

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

- σ_i is a best response to σ_{-i}^i , for every i ;
- σ plays ε -like σ^i , for every i .

Corollary 1

If each σ_i best responds to σ_{-i}^i and $\sigma \ll \sigma^i$ for all i , then for a.e. path $h \exists T$ s.t. $\forall t \geq T$, the continuation σ_{h_t} is a subjective ε -equilibrium.

Proof Idea

Fix $\varepsilon > 0$; for μ^σ -a.e. $h \exists T$ s.t. $\forall t \geq T$, σ_{h_t} plays ε -like $\sigma_{h_t}^i$ for each i (Theorem 1).

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

- σ_i is a best response to σ_{-i}^i , for every i ;
- σ plays ε -like σ^i , for every i .

Corollary 1

If each σ_i best responds to σ_{-i}^i and $\sigma \ll \sigma^i$ for all i , then for a.e. path $h \exists T$ s.t. $\forall t \geq T$, the continuation σ_{h_t} is a subjective ε -equilibrium.

Proof Idea

Fix $\varepsilon > 0$; for μ^σ -a.e. $h \exists T$ s.t. $\forall t \geq T$, σ_{h_t} plays ε -like $\sigma_{h_t}^i$ for each i (Theorem 1).

By rationality, at every t player i plays a best response to $\mu_i(\cdot \mid h_t)$.

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

- σ_i is a best response to σ_{-i}^i , for every i ;
- σ plays ε -like σ^i , for every i .

Corollary 1

If each σ_i best responds to σ_{-i}^i and $\sigma \ll \sigma^i$ for all i , then for a.e. path $h \exists T$ s.t. $\forall t \geq T$, the continuation σ_{h_t} is a subjective ε -equilibrium.

Proof Idea

Fix $\varepsilon > 0$; for μ^σ -a.e. $h \exists T$ s.t. $\forall t \geq T$, σ_{h_t} plays ε -like $\sigma_{h_t}^i$ for each i (Theorem 1).

By rationality, at every t player i plays a best response to $\mu_i(\cdot \mid h_t)$.

Merging $\implies$ those best responses are ε -best responses to true continuation $\mu^\sigma(\cdot \mid h_t)$.

Subjective ε -Equilibrium

Definition (Subjective ε -equilibrium)

A profile $\sigma = (\sigma_i)_i$ is a **subjective ε -equilibrium** if there exist beliefs $\sigma^i = (\sigma_i, \sigma_{-i}^i)$ with:

- σ_i is a best response to σ_{-i}^i , for every i ;
- σ plays ε -like σ^i , for every i .

Corollary 1

If each σ_i best responds to σ_{-i}^i and $\sigma \ll \sigma^i$ for all i , then for a.e. path $h \exists T$ s.t. $\forall t \geq T$, the continuation σ_{h_t} is a subjective ε -equilibrium.

Proof Idea

Fix $\varepsilon > 0$; for μ^σ -a.e. $h \exists T$ s.t. $\forall t \geq T$, σ_{h_t} plays ε -like $\sigma_{h_t}^i$ for each i (Theorem 1).

By rationality, at every t player i plays a best response to $\mu_i(\cdot \mid h_t)$.

Merging $\implies$ those best responses are ε -best responses to true continuation $\mu^\sigma(\cdot \mid h_t)$.

Both (supporting beliefs & closeness) $\implies$ subjective ε -equilibrium from time T .

From Subjective to (Approximate) Nash

Proposition 1

For every $\varepsilon > 0$, $\exists \eta > 0$: if σ is a subjective η -equilibrium then $\exists \sigma^*$ s.t.

- (i) σ plays ε -like σ^* ;
- (ii) σ^* is an ε -Nash equilibrium of the repeated game.

From Subjective to (Approximate) Nash

Proposition 1

For every $\varepsilon > 0$, $\exists \eta > 0$: if σ is a subjective η -equilibrium then $\exists \sigma^*$ s.t.

- (i) σ plays ε -like σ^* ;
- (ii) σ^* is an ε -Nash equilibrium of the repeated game.

Idea: under perfect monitoring and known own payoffs, adjust off-path prescriptions to align incentives while preserving realisations up to ε .

From Subjective to (Approximate) Nash

Proposition 1

For every $\varepsilon > 0$, $\exists \eta > 0$: if σ is a subjective η -equilibrium then $\exists \sigma^*$ s.t.

- (i) σ plays ε -like σ^* ;
- (ii) σ^* is an ε -Nash equilibrium of the repeated game.

Idea: under perfect monitoring and known own payoffs, adjust off-path prescriptions to align incentives while preserving realisations up to ε .

Proof Idea

Fix $\eta > 0$ small. Given subjective η -equilibrium σ , modify off-path prescriptions s.t. unilateral deviations trigger responses that keep the deviator's continuation payoff within ε of best-reply payoff.

Perfect monitoring $\implies$ changes leave realisations ε -close.

Resulting σ^* is an ε -best reply for each player: σ^* is an ε -Nash equilibrium; and σ plays ε -like σ^* .

Main Theorem: Rational Learning $\implies$ Nash Play

Theorem 2 (Kalai and Lehrer 1993)

Suppose each σ_i best responds to σ_{-i}^i and $\mu^\sigma \ll \mu^{\sigma^i}$ for all i . Then for every $\varepsilon > 0$ and for μ^σ -a.e. path h , $\exists T$ s.t. $\forall t \geq T$ there is an ε -Nash equilibrium σ^ε of the repeated game with σ_{h_t} playing ε -like σ^ε .

Main Theorem: Rational Learning $\implies$ Nash Play

Theorem 2 (Kalai and Lehrer 1993)

Suppose each σ_i best responds to σ_{-i}^i and $\mu^\sigma \ll \mu^{\sigma^i}$ for all i . Then for every $\varepsilon > 0$ and for μ^σ -a.e. path h , $\exists T$ s.t. $\forall t \geq T$ there is an ε -Nash equilibrium σ^ε of the repeated game with σ_{h_t} playing ε -like σ^ε .

Proof Idea

- 1) Theorem 1 $\implies$ eventually correct forecasts (merging).
- 2) Best responses to beliefs $\implies \varepsilon$ -best responses to truth (large t).
- 3) Proposition 1 $\implies$ approximate Nash play along the realised path.

Absolute Continuity and Bayesian Nash Equilibrium

Bayesian Nash equilibrium (BNE): in incomplete information (finite type space), each σ_i maximises expected utility given beliefs over types and strategies.

At a BNE of the repeated game, priors give a *grain of truth*: realised play has positive probability under beliefs $\implies$ absolute continuity holds.

Application: starting from a BNE, players eventually play (approximately) a Nash equilibrium of the *realised* complete-information repeated game.

Meaning and Interpretation

What converges? Not actions each period, but *forecasts* of future play; behaviour is best response to (nearly) correct forecasts.

Why it matters: ensures long-run play consistent with Nash discipline without common knowledge of rationality or equilibrium selection.

Learning vs commitment: players learn the environment they *face* (others' strategies), not a fixed state of nature.

Role of absolute continuity: bans dogmatic zero-probability beliefs about realised events; makes Bayes informative.

Learning: with merging, each player's beliefs about future play match the truth; subjective ϵ -equilibrium obtains on-path.

Bayesian Nash starting point

In a repeated game with finitely many payoff types, if play starts at a **Bayesian Nash equilibrium**, then eventually players play (approximately) a Nash equilibrium of the *realised* complete-information repeated game.

Grain of truth at BNE $\implies$ abs. cont.; merging $\implies$ correct forecasts; best responses $\implies$ near-NE of realised environment.

Fudenberg and Levine (1998; 2009 ARE): Main Critique

- Endogeneity of absolute continuity:** abs. cont. must hold for the realised path under the true play; ensuring this is itself an equilibrium-like fixed-point problem.
 - Grain of truth:** wanting priors that *always* put positive mass on the truth is impossible in rich (uncountable) environments; workable classes may be very restrictive.
 - Interpretation caution:** Kalai and Lehrer (1993 Ecta) shows a *consistency* result conditional on abs. cont.; not a general path-to-equilibrium selection theory.
 - Comparative statics:** results sensitive to prior support assumptions; small changes can break abs. cont. and merging conclusion.
 - Bottom line:** powerful when abs. cont. holds (e.g., BNE start with finite types), but limited as a general behavioural foundation without specifying priors.
- “Our interest here, however, is in “learning models,” by which we mean that the allowed priors are exogenously specified, without reference to a fixed point problem.”
Fudenberg and Levine (1998)

Takeaways

Under absolute continuity, Bayesian learning merges beliefs with the truth along realised play.

Rational (best-reply) control with merged beliefs $\implies$ eventual (approximate) Nash play.

At BNE with finite types, eventual play tracks an NE of the realised complete-information game.

Abs. cont. is strong and endogenous; use with care as general foundation for learning in games.

Approachability, Calibration, Adaptive Algorithms, and Sophisticated Learning

Duarte Gonçalves

University College London

Topics in Economic Theory