

Counterfactual Regret Minimization (CFR)

Game Theory · Module 9 — Learning In Games

Node 9.3

Position in Knowledge Graph

System: Game Theory **Structural Domain:** Learning in Games **Node:** 9.3

Counterfactual regret minimization is the algorithm that solved poker. Not approximately, not heuristically — *solved*, in the game-theoretic sense of computing an epsilon-Nash equilibrium for games with 10^{13} to 10^{161} information sets. CFR, introduced by Zinkevich, Johanson, Bowling, and Piccione in 2007, is the structural bridge between the abstract no-regret learning framework of 9.2 and the concrete, astronomically large extensive-form games that define real poker. Its key insight is a decomposition theorem: the total regret of a strategy in an extensive-form game can be written as a sum of *counterfactual regrets* at individual information sets, and if each information set independently minimizes its own counterfactual regret (e.g., via regret matching), the global average strategy converges to a Nash equilibrium of the full game. This decomposition is the structural move that makes enormous games tractable — instead of one monstrous optimization, you get millions of tiny independent ones, each over a handful of actions at a single decision point.

Every major poker AI milestone — Cepheus (2015), DeepStack (2017), Libratus (2017), Pluribus (2019) — is built on CFR or one of its descendants. The algorithm is to poker what the simplex method is to linear programming or what backpropagation is to neural networks: the canonical computational engine that makes the theory useful. This node develops CFR from its formal foundations through its variants (CFR+, Monte Carlo CFR) to its landmark applications, and explains why it is *the* algorithm for imperfect-information games.

Formal Definition

Extensive-form game notation

Let Γ be a finite two-player zero-sum extensive-form game with imperfect information. Define:

- H : the set of all histories (sequences of actions from the root).
- $Z \subseteq H$: the set of terminal histories.
- $A(h)$: the set of actions available at non-terminal history h .
- $P(h) \in \{1, 2, c\}$: the player to act at h (or chance c).
- I_i : the set of information sets of player i . Each $I \in I_i$ is a set of histories that player i cannot distinguish.
- $A(I)$: the actions available at information set I (well-defined because all $h \in I$ have the same action set).
- $u_i(z)$: player i 's payoff at terminal history z .

A **behavioral strategy** for player i is a function σ_i that assigns to each $I \in I_i$ a probability distribution over $A(I)$. A strategy profile $\sigma = (\sigma_1, \sigma_2)$ induces a probability distribution over terminal histories, with expected payoff $u_i(\sigma) = \sum_{z \in Z} \pi^\sigma(z) u_i(z)$, where $\pi^\sigma(z)$ is the probability of reaching z under σ .

Reach probabilities

For a history h , the **reach probability** under strategy profile σ decomposes as:

$$\pi^\sigma(h) = \pi_1^\sigma(h) \cdot \pi_2^\sigma(h) \cdot \pi_c(h),$$

where $\pi_i^\sigma(h)$ is the product of player i 's action probabilities along the path to h , and $\pi_c(h)$ is the product of chance probabilities. Define the **counterfactual reach probability** for player i at h :

$$\pi_{-i}^\sigma(h) = \pi^\sigma(h) / \pi_i^\sigma(h),$$

the probability of reaching h given that player i plays to reach h (i.e., all factors except player i 's own choices).

Counterfactual value

The **counterfactual value** of an information set $I \in I_i$ under strategy profile σ is:

$$v_i^\sigma(I) = \sum_{h \in I} \pi_{-i}^\sigma(h) \sum_{z \in Z(h)} \pi^\sigma(h, z) u_i(z),$$

where $Z(h)$ is the set of terminal histories extending h , and $\pi^\sigma(h, z)$ is the probability of going from h to z under σ . Intuitively: weight each history in the info set by how likely the *opponents and chance* made it, then compute expected payoff under the continuation strategy.

The counterfactual value of a specific action $a \in A(I)$ is:

$$v_i^\sigma(I, a) = \sum_{h \in I} \pi_{-i}^\sigma(h) \sum_{z \in Z(h, a)} \pi^\sigma(h \cdot a, z) u_i(z),$$

where $Z(h, a)$ is the set of terminal histories through h taking action a .

Counterfactual regret

The **instantaneous counterfactual regret** at information set I for action a at iteration t is:

$$r_i^t(I, a) = v_i^{\sigma^t}(I, a) - v_i^{\sigma^t}(I).$$

The **cumulative counterfactual regret** through iteration T is:

$$R_i^T(I, a) = \sum_{t=1}^T r_i^t(I, a).$$

The CFR algorithm

For each iteration $t = 1, 2, \dots, T$:

1. **Traverse the game tree** under current strategy profile σ^t . Compute counterfactual values $v_i^{\sigma^t}(I)$ and $v_i^{\sigma^t}(I, a)$ for every info set I of player i and every action $a \in A(I)$.
2. **Update cumulative regrets:** $R_i^t(I, a) \leftarrow R_i^{t-1}(I, a) + r_i^t(I, a)$.
3. **Compute next strategy via regret matching:**

$$\sigma_i^{t+1}(I)(a) = \begin{cases} \frac{R_i^{t,+}(I, a)}{\sum_{a'} R_i^{t,+}(I, a')} & \text{if } \sum_{a'} R_i^{t,+}(I, a') > 0, \\ \frac{1}{|A(I)|} & \text{otherwise,} \end{cases}$$

where $R_i^{t,+}(I, a) = \max(R_i^t(I, a), 0)$.

4. Accumulate the average strategy:

$$\bar{\sigma}_i^T(I)(a) = \frac{\sum_{t=1}^T \pi_i^{\sigma^t}(I) \cdot \sigma_i^t(I)(a)}{\sum_{t=1}^T \pi_i^{\sigma^t}(I)}.$$

Output: The average strategy profile $\bar{\sigma}^T = (\bar{\sigma}_1^T, \bar{\sigma}_2^T)$.

Intuition

CFR is best understood as an assembly of tiny independent no-regret learners, one per information set, coordinated by the game tree.

The decomposition trick. In a normal-form game, regret matching (9.2) works on a single player’s action set. In an extensive-form game, the “action set” is the space of all behavioral strategies — exponentially large. CFR’s core insight is that global regret *decomposes* into a sum of counterfactual regrets at individual info sets. If each info set independently achieves no-regret (via regret matching), the global strategy achieves no-regret, and by the fundamental theorem of 9.2, the average strategy converges to Nash in zero-sum games.

Counterfactual values as the right weighting. The “counterfactual” in CFR refers to the fact that each info set’s regret is computed *as if the player had tried to reach that info set*. The opponents’ and chance’s contributions to reach probability are included, but the player’s own reach probability is factored out. This is the correct weighting because it ensures that the decomposition of global regret into per-info-set regrets is exact — a non-trivial algebraic identity that Zinkevich et al. proved.

Why “average” strategy, not “current” strategy. At any given iteration, the current strategy σ^t can be wild — it swings from one extreme to another as regrets accumulate and sign-change. The *average* strategy, weighted by reach probabilities, smooths out these oscillations and is the one that converges to Nash. This is analogous to the distinction between last-iterate and average-iterate convergence in optimization: the average iterates of gradient descent converge even when the last iterates cycle.

Why it works for poker. A heads-up no-limit hold’em game has on the order of 10^{161} game states and 10^{14} information sets (after bucketing). No algorithm can enumerate these states explicitly. CFR’s per-info-set decomposition means you only need to store regret values for each (info set, action) pair, not for the global strategy space. With abstraction (clustering similar hands and bet sizes), the number of info sets drops to 10^{12} or fewer, which fits in memory. Each iteration traverses the abstracted tree once, updating regrets. After billions of iterations, the average strategy is an epsilon-Nash equilibrium of the abstract game.

Structural Role

9.3 is the keystone node of the entire game theory course from a poker standpoint. It is where abstract theory meets the concrete reality of billion-dollar poker competition:

- **From 9.2:** CFR inherits the no-regret guarantee and the CCE convergence theorem. The only new ingredient is the counterfactual regret decomposition that lifts normal-form regret matching to extensive-form games.

- **From 3.2 (information sets):** CFR operates on info sets, not game nodes. The entire apparatus of imperfect information — what you know, what you do not, what your strategy must be consistent across — is the structural substrate.
 - **From 3.4 (backward induction):** CFR traverses the game tree recursively, computing values bottom-up like backward induction. The difference is that CFR handles imperfect information and produces mixed strategies rather than pure ones.
 - **From 4.3 (Bayesian games):** The counterfactual reach probability is a belief-like object — “how likely is the opponent to have put me in this info set?” — linking CFR to Bayesian reasoning.
 - **To 9.5 (convergence):** CFR is the constructive proof that learning dynamics converge to Nash in extensive-form zero-sum games.
 - **To 10.3 (strategic systems):** CFR-based solvers are the computational tools that make GTO poker strategies real, not merely theoretical.
-

Mathematical Development

The counterfactual regret decomposition theorem

Theorem (Zinkevich et al. 2007). For any player i in a finite extensive-form game, the overall external regret of player i through T iterations is bounded by the sum of positive counterfactual regrets at each information set:

$$R_i^T \leq \sum_{I \in I_i} R_i^{T,+}(I),$$

where $R_i^{T,+}(I) = \max_{a \in A(I)} R_i^T(I, a)$ is the maximum positive counterfactual regret at I , and R_i^T is the overall external regret.

Consequence. If each info set I achieves sublinear counterfactual regret — i.e., $R_i^{T,+}(I) = O(\sqrt{T})$ — then the overall regret is $R_i^T \leq |I_i| \cdot O(\sqrt{T}) = O(|I_i| \sqrt{T})$. The average strategy’s exploitability is:

$$\epsilon^T \leq \frac{R_1^T + R_2^T}{T} = O\left(\frac{|I_1| + |I_2|}{\sqrt{T}}\right).$$

To achieve exploitability ϵ , you need $T = O((|I_1| + |I_2|)^2 / \epsilon^2)$ iterations. For a game with 10^{12} info sets and target exploitability 10^{-3} of the pot, this is astronomical — which is where Monte Carlo CFR and abstraction enter.

Proof sketch of the decomposition

The key identity: for any pair of strategies σ, σ' for player i , the difference in expected payoff can be written as:

$$u_i(\sigma'_i, \sigma_{-i}) - u_i(\sigma_i, \sigma_{-i}) = \sum_{I \in I_i} \pi_{-i}^\sigma(I) \sum_{a \in A(I)} [\sigma'_i(I)(a) - \sigma_i(I)(a)] v_i^\sigma(I, a).$$

This follows from the factored form of reach probabilities and the law of total expectation over the game tree. The counterfactual regret of not playing action a at info set I is exactly the per-info-set contribution to the global regret of not switching to a everywhere at I . The decomposition is exact, not an approximation.

CFR+ (Tammelin 2014)

CFR+ is a practical improvement that accelerates convergence significantly:

1. **Regret flooring:** Instead of allowing cumulative regrets to go arbitrarily negative, CFR+ floors them at zero: $R_i^t(I, a) \leftarrow \max(R_i^{t-1}(I, a) + r_i^t(I, a), 0)$. This prevents an action that was bad early on from accumulating a large negative regret that takes many iterations to “pay off” before the action can be reconsidered.
2. **Linear averaging:** Instead of uniform averaging of past strategies, CFR+ weights iteration t by t in the average strategy:

$$\bar{\sigma}_i^T(I)(a) \propto \sum_{t=1}^T t \cdot \pi_i^{\sigma^t}(I) \cdot \sigma_i^t(I)(a).$$

This gives more weight to later iterations, which are closer to equilibrium, and speeds up convergence in practice.

Empirical performance. CFR+ converges roughly 10x faster than vanilla CFR in heads-up limit hold’em, and was the algorithm used to solve that game (Cepheus, Bowling et al. 2015). The theoretical convergence guarantee is the same — $O(T^{-1/2})$ exploitability — but the constants are dramatically better.

Monte Carlo CFR (MCCFR)

The bottleneck of vanilla CFR is that each iteration requires a full traversal of the game tree, computing counterfactual values at every info set. For large games (HUNL has $\sim 10^{161}$ game states), full traversal is impossible. Monte Carlo CFR replaces full traversal with *sampled* traversals, updating only the info sets encountered on a sampled path through the tree.

External sampling. Sample the opponent’s actions and chance outcomes; traverse the tree fully for the updating player’s actions. At each iteration, only the info sets along the sampled path are updated. This reduces per-iteration cost from $O(|H|)$ to $O(|H_i|)$ (the tree restricted to player i ’s decisions with opponent/chance actions sampled).

Outcome sampling. Sample *all* actions — both players and chance — to get a single trajectory through the tree. Update regrets only along that trajectory. Per-iteration cost is $O(d)$ where d is the depth of the tree, but variance is high.

Chance sampling. Sample only chance outcomes (e.g., which cards are dealt), then traverse the resulting perfect-information game fully. This is the most common variant for poker, because chance sampling eliminates the card-dealing combinatorics (which dominate the branching factor) while still doing exact computation within each deal.

Convergence. All MCCFR variants converge to Nash in the same $O(T^{-1/2})$ sense, but with higher variance. The effective convergence speed depends on the sampling scheme’s variance-to-cost tradeoff. In practice, chance sampling with pruning is the standard for poker applications.

Pruning

Regret-based pruning (Brown and Sandholm 2015). If all actions at an info set have cumulative regret below $-C$ for some threshold, skip that info set entirely for the next several iterations — the regrets cannot become positive quickly, so the strategy will not change. This can reduce per-iteration cost by 10-100x in practice without affecting convergence.

Abstraction

Real poker games are too large for CFR even with Monte Carlo sampling. **Abstraction** reduces the game to a tractable size:

- **Card abstraction:** Cluster similar hands into “buckets” based on equity distribution, hand strength, or learned features. E.g., all flush draws on a monotone flop might be one bucket.
- **Action abstraction:** Restrict bet sizes to a finite grid (e.g., 33%, 67%, 100%, 150% of pot). The continuous bet-size space is discretized.

CFR is run on the abstract game; the resulting strategy is an approximate equilibrium of the abstract game. Whether it is a good approximation to the *real* game depends on the quality of the abstraction. Pathological abstractions can introduce significant error (the “abstraction pathology” — Waugh et al. 2009).

Subgame solving

Rather than solving the entire game once, decompose the game into subgames and solve each independently:

- **Unsafe subgame solving:** Fix the strategy in the rest of the tree and solve a subgame in isolation. This can produce exploitable strategies because the subgame solution may not be consistent with the trunk strategy.
- **Safe subgame solving (Burch, Johanson, Bowling 2014; Brown, Sandholm 2017):** Maintain a “gift” or “floor” value for each subgame entry point, ensuring that the subgame solution does not make the overall strategy worse. Libratus and Pluribus use safe subgame solving to resolve postflop situations in real time during play.

Depth-limited solving. Instead of solving all the way to terminal nodes, solve to a limited depth and use a learned value function to estimate continuation values at the leaves of the subtree. DeepStack uses this approach with neural-network value estimation.

Worked Example

CFR on Kuhn Poker

Kuhn Poker is the canonical toy poker game for demonstrating CFR. Three cards: Jack (J), Queen (Q), King (K). Two players. Each antes 1 chip. One card dealt to each player. Player 1 acts: check or bet 1. If check, Player 2 acts: check or bet 1. If Player 2 bets after check, Player 1 acts: fold or call. If Player 1 bets initially, Player 2 acts: fold or call. Higher card wins at showdown.

Information sets. Player 1 has 6 info sets: $\{J, Q, K\} \times \{\text{first action, facing bet after check}\}$. Player 2 has 6 info sets: $\{J, Q, K\} \times \{\text{facing check, facing bet}\}$.

CFR iteration 1. Initialize all strategies to uniform. Traverse the tree for all 6 possible card deals (J-Q, J-K, Q-J, Q-K, K-J, K-Q, each equally likely). For each deal, compute counterfactual values at each info set using the uniform strategy. Update cumulative regrets. Compute new strategy via regret matching.

After a few hundred iterations, the average strategy converges to the known Nash equilibrium of Kuhn Poker:

Player 1 (Nash equilibrium): - Holding J: check 100%. If facing bet: fold 2/3, call 1/3. - Holding Q: check 100%. If facing bet: call 1/3, fold 2/3. - Holding K: bet $3 \times \alpha$ for parameter $\alpha \in [0, 1/3]$; check otherwise. If facing bet after check: call 100%.

Player 2 (Nash equilibrium): - Holding J, facing check: bet 1/3. Facing bet: fold 100%. - Holding Q, facing check: check 100%. Facing bet: call 1/3, fold 2/3. - Holding K: always bet or call.

The game value is $-1/18$ for Player 1 (Player 2 has a slight positional advantage). CFR reaches exploitability below 10^{-4} within a few thousand iterations on this toy game.

Exploitability decay in Kuhn Poker

Iterations	Exploitability (chips)
10	0.15

Iterations	Exploitability (chips)
100	0.04
1,000	0.012
10,000	0.003

The decay tracks the theoretical $O(T^{-1/2})$ rate. CFR+ on the same game converges roughly 5x faster.

Poker Anchor

Concept mapping

Formal object	Poker instantiation
Information set I	A specific decision point: your hand, the board, the action history so far
Action set $A(I)$	Fold, call, raise (various sizes) at that decision point
Counterfactual value $v_i^\sigma(I)$	The EV of reaching this spot, weighted by opponents' and chance's contribution to getting here
Counterfactual regret $r_i^t(I, a)$	How much better action a would have been than the current mixed strategy at this info set, this iteration
Cumulative regret $R_i^T(I, a)$	The running total of how much action a has been "missed" — positive means the action is underplayed
Regret matching	Mix over actions in proportion to how much you "regret" not playing them — the more you regret not raising, the more you raise next iteration
Average strategy $\bar{\sigma}^T$	The GTO output: the strategy the solver reports after T iterations
Exploitability	How much an opponent can win by best-responding to your strategy — the solver's error metric
Abstraction	Bucketing hands and discretizing bet sizes to make the game tree tractable
Subgame solving	"Solve this river spot in real time" — what Libratus and Pluribus do during play

Canonical poker applications

Cepheus (2015, Bowling et al., University of Alberta). Heads-up limit Texas hold'em (HULHE) was "essentially solved" using CFR+. The game has approximately 3.19×10^{14} information sets after isomorphic reduction and 3.16×10^{17} game states. Bowling et al. ran CFR+ for over 1,500 CPU-years of computation, producing a strategy with exploitability less than 1 milli-big-blind per hand — so small that a human playing 60 hands per hour for a 70-year lifetime would not be able to statistically detect the deviation from Nash play. Cepheus was the first *non-trivial* game to be "essentially weakly solved" in the sense of computational game theory: the strategy's exploitability is within the noise floor of statistical significance.

Libratus (2017, Brown and Sandholm, CMU). Libratus defeated four top professionals in a 120,000-hand heads-up no-limit Texas hold'em (HUNL) match ("Brains vs. Artificial Intelligence"), winning by a combined 1.77 million chips (14.7 bb/100 — a decisive margin). Libratus used three modules:

1. **Blueprint computation:** An abstracted HUNL game with $\sim 10^{12}$ info sets solved by Monte Carlo CFR to produce a coarse strategy for the entire game tree.
2. **Subgame solving:** When a specific postflop situation was reached, Libratus re-solved the relevant subgame in real time using safe subgame solving with a finer action abstraction. This is the move

that made HUNL tractable: the blueprint handles the macro-structure, and real-time subgame solving provides precision at the micro-level.

3. **Self-improvement:** After each day of play, the blueprint was refined to close exploits the human opponents had found, using a targeted regret-update procedure.

Libratus’s victory was the first time an AI beat top human players in HUNL, a game with $\sim 10^{161}$ game states.

DeepStack (2017, Moravcik et al.). DeepStack took a different approach: instead of computing a full blueprint, it used *depth-limited solving* with a deep neural network to estimate continuation values beyond the solving horizon. At each decision point, DeepStack re-solves the next few actions using CFR with neural-net leaf values. The neural nets were trained on millions of randomly generated poker situations, providing a fast approximate value function. DeepStack defeated professional poker players in HULHE matches, and its approach was the first to integrate deep learning with game-theoretic solving.

Pluribus (2019, Brown and Sandholm). Pluribus extended the Libratus approach to *six-player* no-limit hold’em — a much harder problem because (a) the game is no longer two-player zero-sum, (b) the number of possible opponent interactions grows combinatorially, and (c) Nash equilibrium is no longer the unique “safe” solution concept (CCE is the natural target). Pluribus used:

1. **Monte Carlo CFR for blueprint computation** in the six-player game, with action abstraction and card abstraction.
2. **Depth-limited search** during play, solving subgames to a limited depth with blueprint-based continuation values.
3. **A modified search that considers correlated opponent responses** — notably, Pluribus does *not* assume opponents play independently, and its real-time search considers opponent action profiles that are correlated in response to Pluribus’s deviations.

Pluribus defeated 12 elite professionals (including several World Series of Poker bracelet winners) over 10,000 hands of 6-max no-limit hold’em, with a win rate statistically significant at $p < 0.001$. Its training cost was approximately \$150 of cloud compute — roughly 8 days on a 64-core server — a striking demonstration that CFR-based methods do not require supercomputer-scale resources.

Why CFR is THE algorithm for poker

Five structural reasons:

1. **Handles imperfect information natively.** Unlike backward induction or minimax search, CFR operates on information sets, not game nodes. It does not need to know the opponent’s cards — it works with the beliefs induced by the info-set structure.
2. **Decomposes over info sets.** The counterfactual regret decomposition means the algorithm scales with the number of *info sets*, not the number of game *states*. In HUNL, the number of game states is 10^{161} , but with abstraction, the number of info sets can be reduced to 10^{12} or fewer.
3. **Guaranteed convergence to Nash in zero-sum games.** The no-regret guarantee at each info set implies that the average strategy converges to a Nash equilibrium. This is not a heuristic — it is a theorem.
4. **Compatible with Monte Carlo sampling.** Full tree traversal is impossible for large games, but MCCFR variants allow sampled traversals that still converge, making CFR practical for games with $10^{14}+$ info sets.
5. **Compatible with subgame solving and depth-limited search.** The blueprint-plus-refinement architecture used by Libratus and Pluribus is a natural extension of CFR: compute a coarse solution offline, then refine specific subgames online. This two-level structure is what makes real-time play feasible.

No other known algorithm family combines all five properties. Alpha-beta search (perfect-information games only), linear programming (does not scale beyond small games), fictitious play (no per-info-set decomposition,

slow convergence), and reinforcement learning (no convergence guarantees in multi-agent settings) each fail on at least one of these axes.

Concrete situation: what your solver is doing

When you open PioSolver and input a flop spot — say, BTN vs BB single-raised pot on $A\heartsuit 7\diamondsuit 2\clubsuit$ — the solver is running CFR (or a variant) on the game tree rooted at that flop node. It initializes uniform strategies at every info set. Each iteration:

1. Traverses the tree from the flop through all turns and rivers (with chance sampling for card deals).
2. At each info set (your hand + board + action history), computes counterfactual values for each action.
3. Updates cumulative regrets.
4. Computes next-iteration strategy via regret matching.
5. Accumulates the average strategy.

After the configured number of iterations (typically 200-2000 for a subgame), the solver reports the average strategy and its exploitability. The output — “bet 67% pot with this range, check with that range” — is the average strategy $\bar{\sigma}^T$, and the exploitability number is $(R_1^T + R_2^T)/T$.

Concrete situation: real-time subgame re-solving at the table

Imagine you are playing a HUNL match and reach the turn on a board that was not in your pre-computed solution’s abstraction grid. A Libratus-style approach would:

1. Take the current reach probabilities (your range and opponent’s range at this turn node, implied by the blueprint strategy).
2. Construct a subgame from this turn node through all rivers to showdown, using a fine action abstraction.
3. Run CFR on this subgame for enough iterations to reach low exploitability.
4. Play according to the resulting strategy for this specific subgame.

This is what makes modern poker AI superhuman: it does not memorize a single fixed strategy. It recomputes, on the fly, the Nash equilibrium of each specific subgame it encounters, using CFR as the computational engine.

What this understanding buys you

Every solver output is a CFR average strategy. When you study a solver solution, you are looking at the output of a no-regret learning process that ran for many iterations. The strategy is not “the answer” — it is the average of thousands of iterations of a self-play process, and its quality is bounded by the exploitability metric. If the exploitability is 1% of pot, the strategy is within 1% of Nash. If it is 10%, you are still far from convergence.

Abstraction quality matters more than iteration count. The biggest source of error in solver output is not “not enough iterations” (which produces a known, bounded error) but “bad abstraction” (which can produce arbitrarily bad strategies). If you bucket hands poorly or discretize bet sizes too coarsely, CFR will converge to the Nash of the *wrong game* — perfectly optimal in the abstract game, but exploitable in the real game. This is why serious solvers spend significant effort on abstraction design.

CFR explains why GTO is hard to compute but easy to approximate. The exact Nash equilibrium of HUNL is intractable (the game is too large). But an epsilon-Nash with $\epsilon = 1$ mbb/hand is “essentially solved.” CFR is the algorithm that makes this possible: it does not find the exact Nash, but it finds a strategy that is so close to Nash that no opponent can exploit the gap in a human lifetime of play.

CFR explains the structure of solver outputs. Why do solver solutions often involve mixing between actions at specific frequencies? Because regret matching at each info set produces a mixed strategy that balances the regrets of each action. The frequencies are not arbitrary — they are the steady-state of a no-regret learning process. When a solver says “bet 67% pot 40% of the time, check 60% of the time,” it is reporting the regret-matching equilibrium at that info set.

Cross-Links

- **9.1 Fictitious play** — the ancestor algorithm; CFR improves on it with per-info-set decomposition and no-regret guarantees.
 - **9.2 No-regret learning** — the theoretical foundation; CFR is regret matching applied to extensive-form games.
 - **9.4 Multiplicative weights** — an alternative no-regret algorithm; could replace regret matching at each info set in CFR (rarely done in practice).
 - **9.5 Convergence to equilibrium** — CFR is the constructive proof that no-regret dynamics converge to Nash in extensive-form zero-sum games.
 - **3.2 Information sets** — the structural substrate on which CFR operates.
 - **3.4 Backward induction** — CFR's tree traversal is structurally analogous, but handles imperfect information.
 - **4.3 Bayesian games** — counterfactual reach probabilities encode beliefs about opponent types.
 - **2.1 Nash equilibrium** — the convergence target.
 - **Online learning / optimization** — regret minimization in the extensive-form setting.
-

Structural Compression

Invariant: In a finite two-player zero-sum extensive-form game, the total external regret decomposes exactly into a sum of counterfactual regrets at individual information sets; independent no-regret learning (regret matching) at each info set therefore drives the average strategy profile to a Nash equilibrium, with exploitability decaying as $O(T^{-1/2})$.

Mechanism: Counterfactual reach probabilities factor out the acting player's contribution, enabling a per-info-set decomposition of global regret. Regret matching at each info set produces a mixed strategy proportional to positive cumulative regrets. The average strategy across iterations is the output. Monte Carlo sampling, pruning, and subgame solving extend the method to games too large for full traversal.

Cross-System Echo: CFR is the game-theoretic analogue of **coordinate descent in optimization**: instead of optimizing a single massive objective, decompose it into independent per-coordinate subproblems and cycle through them. In machine learning, stochastic gradient descent decomposes the empirical risk over data points; in CFR, the decomposition is over information sets. The structural insight is the same — *global convergence from local updates* — and the mathematical engine is the same — *each local update guarantees bounded local regret, and the global bound is the sum*. The decomposition theorem (Zinkevich 2007) is the game-theoretic counterpart of the separability lemma in coordinate descent theory.

Exercises

1. Implement vanilla CFR for Kuhn Poker. Run 10,000 iterations and verify that the average strategy's exploitability is below 0.005 chips. Compare the output to the known analytic Nash equilibrium.
2. Prove the counterfactual regret decomposition theorem: show that $R_i^T \leq \sum_{I \in I_i} R_i^{T,+}(I)$ for any extensive-form game. (Hint: expand the definition of global external regret using the reach-probability factorization.)
3. Derive the exploitability bound $\epsilon^T \leq (|I_1| + |I_2|) \cdot C/\sqrt{T}$ from the per-info-set $O(\sqrt{T})$ regret bound. What does this imply about the number of iterations needed to solve a game with 10^{12} info sets to exploitability 10^{-3} ?

4. Explain why CFR+ (regret flooring + linear averaging) converges faster than vanilla CFR in practice. Construct a simple example (e.g., a modified Kuhn Poker) where an action accumulates large negative regret under vanilla CFR that delays convergence, and show that CFR+ avoids this.
5. Compare external sampling and outcome sampling in MCCFR. For a game tree of depth d with branching factor b , what is the per-iteration cost of each? Which has higher variance? Why is chance sampling the standard in poker?
6. Explain the abstraction pathology: construct a simple extensive-form game where applying a hand abstraction (merging two info sets into one) causes CFR to converge to a strategy that is more exploitable than the original game's Nash equilibrium. What structural property of the abstraction causes this failure?
7. **Argue, structurally**, why CFR's per-info-set decomposition is the essential algorithmic innovation for imperfect-information games, and why no algorithm based on full-game optimization (e.g., linear programming for the sequence form) can scale to games of poker's size. What is the computational complexity barrier that decomposition overcomes?

Game Theory · Module 9 — Learning In Games · Node 9.3

Concepts: convergence, nash-equilibrium, computational-complexity

Prerequisites: 9.1, 9.2, 3.2, 3.4, 4.3

Enables: 9.5, 10.3

hbar.university — own.your.cognition