

# MirageV145: Nash-Anchored Hardball with Adversarial Concealment and a Scenario-Adaptive Opponent Model

Automated Negotiation League (ANL) 2026 — ANAC @ IJCAI 2026 — NegMAS / Python

Author: \_\_\_\_\_ Affiliation / Team: \_\_\_\_\_

**Abstract.** MirageV145 is our entry for the ANL 2026 bilateral, multi-issue negotiation league, whose score rewards both individual *advantage* and *preference concealment* (a Kendall- $\tau$  term comparing each agent's reported model of the opponent to the truth). MirageV145 pairs a hard, Nash-anchored Boulware concession — which dominates the advantage term — with three mechanisms that act *only* on the  $\tau$  half of the score, leaving the advantage-winning bidding byte-for-byte unchanged: (i) it balances the value-frequencies it reveals so a frequency-modelling opponent cannot infer its preferences (lowering  $\tau_{\text{opp}}$ ); (ii) it never reports a flat ranking (no  $\tau$  forfeit); and (iii) — new in this version — it reports a *scenario-adaptive* opponent model that infers, on-line, whether the negotiation is competitive or cooperative and shapes its reported preferences accordingly, correcting a systematic error of fixed competitive priors. The result is an agent that is advantageous, hard to read, and accurate at modelling opponents across the full competitive-cooperative spectrum.

## 1. Setting and Design Principles

An ANL 2026 agent's per-negotiation score is *advantage* +  $\tau/(\tau + \tau_{\text{opp}})$ , where *advantage* =  $(u(\omega) - r) / (\max u - r)$  for an agreement  $\omega$  and reservation value  $r$ , and  $\tau = \frac{1}{2}(1 + \kappa(\hat{u}_B, u_B))$  rewards the rank-accuracy (Kendall  $\tau$ -b,  $\kappa$ ) of the agent's reported model of the opponent's utility. The opponent's true utility is hidden; each side only observes the other's offers. Two design consequences drive MirageV145:

**(1) Every no-agreement scores zero advantage**, so concession must secure a rational deal whenever a zone of agreement exists — yet conceding too early surrenders advantage. We resolve this with a steep (Boulware) concession that holds firm until late, floored at a fraction of the Nash-bargaining outcome so we never accept a poor deal.

**(2) The  $\tau$  term has two levers**: model the opponent accurately (raise  $\tau$ ) and reveal little about oneself (lower the opponent's  $\tau$ ,  $\tau_{\text{opp}}$ ). MirageV145 attacks both, and does so *without* changing the advantage-winning bidding. This separation is deliberate: in repeated live evaluation, every modification of the concession/bidding engine itself — deal-closing, expected-utility optimal-stopping, relative-risk, integrative logrolling, endgame reservation-value targeting — reduced live performance relative to disciplined hardball, whereas changes confined to the reported model and to bid-information-leakage carried no such risk. MirageV145 therefore restricts all of its innovation to the  $\tau$  half.

## 2. Strategy

**Rational set and Nash-anchored floor.** On receiving preferences, MirageV145 enumerates the outcome space (sampling a bounded representative subset for very large domains to keep every step within the per-negotiation time limit), keeps those above its reservation value, and estimates the Nash-bargaining point of the (self, estimated-opponent) pair over a bounded frontier. Its floor is  $\max(r, 0.95 \cdot u_{\text{self}}(\text{Nash}))$ ; it never offers or accepts below this. The opponent's reservation value used in the Nash computation is estimated as the lowest self-utility the opponent has voluntarily offered.

**Concession.** Aspiration follows a steep Boulware curve,  $\text{asp}(t) = \max\_u - (\max\_u - \text{floor}) \cdot t^\beta$  with  $\beta=64$ , holding near the maximum and conceding only toward the deadline. Acceptance meets a fraction of the current aspiration, guarantees any above-reservation deal at the deadline, and never accepts below the reservation value. This engine is identical to our proven hardball baseline; MirageV145 changes nothing here.

**Adversarial concealment (lowering  $\tau_{\text{opp}}$ ).** Among the bids that clear the current aspiration, MirageV145 (a) minimises the rank-correlation that a frequency-modelling opponent would infer from its bid history, and (b) *balances the value-frequencies it reveals* — it avoids repeatedly offering the same values, flattening its revealed value distribution so the opponent's frequency model cannot identify which values MirageV145 prefers. Because this selection acts only among bids that already clear the aspiration, our own utility is unchanged.

**Scenario-adaptive opponent model (raising  $\tau$ , new in v145).** MirageV145 maintains an online hard-headed frequency model of the opponent. The key innovation is in the model it *reports*: rather than break ties with a

fixed *competitive* prior (assuming the opponent's preferences oppose ours) — which is correct in zero-sum scenarios but actively wrong in cooperative ones, silently forfeiting  $\tau$  there — MirageV145 estimates the scenario's character on-line. It measures whether the opponent's revealed offers sit, on average, *above* or *below* our own mean utility: offers that are above-average for us indicate aligned (cooperative) preferences, below-average indicate opposed (competitive) preferences. The reported utility is then a shrinkage blend,  $\alpha \cdot \text{freq} + (1-\alpha) \cdot \text{prior}$ , where  $\alpha$  grows with the number of observed offers (prior-driven when data is scarce, frequency-driven once informed, with a small persistent prior weight to correct systematic frequency error), and the prior interpolates between competitive ( $1-u_{\text{self}}$ ) and cooperative ( $u_{\text{self}}$ ) according to the estimated alignment. A negligible tie-breaker guarantees the ranking is never flat. This adapts our reported model to the true geometry of each scenario without ever consulting the hidden opponent utility.

**Statelessness and safety.** All state lives in instance attributes reset per negotiation; the agent writes nothing to disk, globals, or class-level variables, contains no print statements, and is a single self-contained module — satisfying the competition's rules of participation. The scenario-adaptive estimate is derived solely from offers observed within the current negotiation.

### 3. Experimental Findings

Across local NegMAS cartesian tournaments on randomly generated 2–5 issue scenarios and live ANL leaderboard cycles we observed:

- **Hard concession is the robustly competitive class.** The hardball core held strong leaderboard positions across cycles, whereas softer or adaptive bidding paradigms (deal-closing, expected-utility optimal-stopping, relative-risk, integrative logrolling, endgame targeting) each under-performed in the live field. We therefore confine innovation to the  $\tau$  half and keep the bidding engine fixed.
- **The scenario-adaptive report raises  $\tau$  at no advantage cost.** In head-to-head local evaluation against the fixed-competitive-prior report, the adaptive model raised our mean opponent-model accuracy (Kendall  $\tau$ ) by approximately +0.05 — the largest single-mechanism  $\tau$  gain we have measured — while individual advantage and agreement rate were statistically unchanged, as expected since the bidding is byte-identical. The gain concentrates in cooperative scenarios, where the fixed competitive prior had been near-worst-case.
- **Local advantage is a weak predictor of live rank**, so we validate strategic changes on the  $\tau$  (estimation-accuracy) half, which is scenario-general and transfers, and treat advantage-side changes as live experiments to be avoided unless proven.

### 4. Conclusion

MirageV145 demonstrates that in a hidden-preference, concealment-scored negotiation, a disciplined hard-Boulware core — protected by a Nash-anchored floor and augmented with adversarial preference concealment and a scenario-adaptive, degeneracy-hardened opponent-model report — yields an agent that is both individually advantageous and accurately models opponents across competitive and cooperative scenarios. By confining every innovation to the  $\tau$  half and leaving the advantage-winning bidding intact, MirageV145 is a strict, low-risk improvement over the unconcealed hardball baseline.