



Comparative Analysis of Five Global Negotiation Frameworks Using AI-Based Simulation: A Matched-Pair Experimental Design

Keld Jensen

Business Administration / Management (Negotiation & Strategy).

Article Information

Received: September 10, 2026

Accepted: September 18, 2026

Published: September 21, 2026

Corresponding author: Keld Jensen, Business Administration / Management (Negotiation & Strategy).

Citation: Jensen K., (2026) "Comparative Analysis of Five Global Negotiation Frameworks Using AI-Based Simulation: A Matched-Pair Experimental Design" Journal of Social and Behavioral Sciences, 4(1); DOI: 10.61148/3065-6990/JSBS/072.

Copyright: ©2026. Keld Jensen. This is an open access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

Abstract

A rule-based simulation of 4,000 agents was performed in an industrial procurement scenario to compare 5 negotiation frameworks. The models that were tested include the SMARTnership model (Jensen), the Harvard Principled Negotiation model (Fisher and Ury), the Structured Model of Giuseppe Conti, The Gap Partnership framework, and the Black Swan Method of Chris Voss. Each buyer-supplier pair negotiated under all five frameworks, resulting in an 800 buyer-supplier pairs within-subjects design. Agents negotiated over four interdependent issues in 10 rounds.

Three outcome variables were used to measure performance: joint value captured, concession-process quality and outcome asymmetry. A fourth variable, the trust index, remained stable within four of the five frameworks, and is presented as a characteristic of the agent architecture at design level and not as a dependent variable.

Friedman test results showed that there were significant differences among frameworks for all three outcome variables ($p < 0.001$). The SMARTnership agent achieved the highest joint value (93.78% of Pareto-efficient potential; paired Cohen's $d = 0.99 - 1.99$ vs. comparison frameworks). The Harvard agent had the highest score in the smoothness of the concession process and reciprocity (9.69 on a 0 to 10 scale), and the lowest score in the asymmetry of the value distribution (asymmetry index 0.040). There was no significant difference between SMARTnership and Harvard for asymmetry (adjusted $p = 1.00$). The Voss agent was dead last on all three measures. The Conti and Gap Partnership agents sat in a comfortable middle ground.

There was no dominant framework on all three dimensions. The agents using transparency were the most successful in increasing the surplus available, the interest-based agent in splitting the surplus most fairly and the smoothest concession path. The results presented here provide a description of the modeled system's behavior and are not directly transferable to trained human negotiators. agent-based negotiation, matched-pair simulation, joint value, concession-process quality, outcome asymmetry, SMARTnership, Harvard Principled Negotiation

Keywords: agent-based negotiation; matched-pair simulation; joint value; concession-process quality; outcome asymmetry; SMARTnership; Harvard Principled Negotiation

1. Introduction

Negotiation sits at the centre of how organisations allocate resources, resolve disputes, and build Negotiation is at the heart of the ways organisations allocate resources, solve problems and forge partnerships. In the last few decades, a few models have emerged that have defined practitioner training and academic debate: the interest-based approach of Fisher and Ury (1981); the tactical psychology of Voss and Raz (2016); the structured rationalism of Conti (2019); the behavioural process discipline of The Gap Partnership (2023); and the transparency-based economics of Jensen (2018). Although these frameworks differ in their underlying principles, few studies have compared their performance directly under the same conditions. Much of the evidence is from case studies, small group experiments and practice reports, which are limited in size, variation of situation and difficulty in controlling for negotiator skill.

Computational simulation is a solution to these limitations. By coding the logic of each framework into separate agents, and by conducting thousands of negotiations from the same initial conditions, the effect of the framework can be separated from the ability of the agent implementing it. However, the agents are simplified representations of human negotiators and cannot capture factors such as experience, judgement, or interpersonal behaviour. In this study, the five frameworks are compared by using a matched-pair agent simulation. Eight hundred buyer-supplier dyads with randomly assigned endowments negotiated in all five of the frameworks resulting in 4,000 negotiations with full logs of offers at each level of negotiation. The matched design allows performance differences between frameworks to be examined under the same negotiation conditions.

Three outcome variables are analyzed: joint value captured, which indicates the amount of the available surplus that the agents captured; concession-process quality, which indicates the smoothness and reciprocity of the bargaining path; and outcome asymmetry, which indicates how evenly the final agreement split the value between buyer and seller. The fourth variable, the trust index, was found to be constant within four of the five frameworks and is considered to be an architectural variable and not an outcome. The study therefore examines whether the frameworks perform differently across the three outcome dimensions.

2. Literature Review

2.1 Negotiation Frameworks in Theory and Practice

Research on negotiation has developed across psychological, strategic, and behavioural perspectives. Early scholars treated negotiation mainly as a psychological issue related to persuasive, concession-making and perception management in the face of uncertainty (Spector, 1977). Hausken (1997) added strategic rationality to this scene and cast the negotiator not as an emotional player but as a utility-maximizing agent making cooperative or competitive moves. For practical applicability, these formal models were less important than for what they had created conceptually: negotiation outcomes were more a function of information, incentives, and decision structure than of personality. The most significant practical change was the introduction of interest-based bargaining by Fisher and Ury (1981) who made it a trainable methodology. The Harvard model suggested four disciplines: separate the people from the problem, concentrate on interests, not positions, invent options for mutual gain, and commit to objective standards. The practical value was immediately

apparent, and the model prevailed in legal, diplomatic and commercial education for the next forty years (Druckman, 2010). What it did not do, however, was to give a systematic account of how preference disclosure impacts the set of attainable outcomes, which Jensen (2018) subsequently did.

The Black Swan Method of Voss and Raz (2016) is from a different tradition altogether. While Fisher and Ury were interested in mutual satisfaction, Voss was interested in information advantage, using calibrated empathy, emotional labelling and the strategic use of silence to shape the opponent's psychological state before substantive bargaining. Schweitzer, DeChurch, and Gibson (2005) had previously shown that competition affects the offers that negotiators make as well as the type of deceptive tactics they will use. The Black Swan Method incorporates these competitive dynamics into a practical negotiation approach. Some say it takes value at a relational expense which only shows up in subsequent transactions; others say that in one-shot, high-stakes transactions, it just works better than principled goodwill.

Conti (2019) and The Gap Partnership (2023) fall somewhere in the middle of these extremes, but with different foci. Conti emphasizes the importance of preparation discipline and the logical sequence of trade-offs, and negotiators enter each session with a ranked list of concessions, a clear reservation point, and a resistance strategy based on the opposition they expect. The Gap Partnership takes the rationalist approach and introduces behavioural situational awareness with conditionally concession logic and procedural checkpoints. Neither maximise joint gains like Harvard or SMARTnership, and both claim to deliver repeatable, auditable outcomes, which is a different value proposition and one particularly relevant for procurement environments where repeatability of the process is more important than creative surplus.

Jensen (2018) builds on the Harvard tradition, making transparency itself an economic variable. Logrolling is possible when both sides reveal their actual priorities of the issues, because by giving up on low-priority issues they can gain on high-priority ones, and the resulting trades increase the total surplus to be divided. This is based on the negotiation analysis literature (Lax & Sebenius, 1986; Raiffa, 1982), which showed that Pareto improvements can be made that are not available to positional bargainers when preference heterogeneity is known. Jensen (2018) proposed a disclosure-based methodology for adversarial commercial negotiations, where revealing true preferences involves considerable strategic risk.

The links between these frameworks and outcomes have been well theorised but less tested comparatively. Druckman (2010) pointed this out explicitly: Negotiation scholarship had created sophisticated models, but there was a lack of controlled comparative evidence that would enable practitioners to select between them on the basis of results. The present study is a response to that observation.

2.2 What Computational Simulation Adds

It is not the case that computational negotiation simulation is a matter of agents behaving like people. They do not. The reason is that there is no other controlled comparison under the same circumstances. If two trained practitioners negotiate, each with a different framework, the result is a product of their frameworks, their skill, their history with each other, their cognitive state on the day, and many other factors that cannot be controlled. All of these

are removed by a computational agent except the framework.

Initial automated negotiation research provided the basic framework of agents with utility functions that make offers and concession schedules and belief updating rules that are passed around a fixed number of times (Debenham, 2003, 2004). Sierra, Jennings, Noriega, and Parsons (1997) and Amgoud, Dimopoulos, and Moraitis (2007) incorporated argumentation-based mechanisms that enable agents to provide justification for proposals, based on structured reasoning, instead of just utility optimisation. This was later broadened to adaptive systems that could change their strategy in the middle of negotiation based on observed opponent behaviour by Zhang, Shen, and Ghenniwa (2008).

One thing that was always missing from this literature was any linkage to named practitioner frameworks. The agents were designed to test abstract propositions, e.g., early concession vs. late concession, preference disclosure vs. no preference disclosure, but not to map these propositions onto the specific set of rules that practitioners learn. Lai and Sycara (2009) were the closest to designing a generic multi-attribute architecture that can be flexible enough to encode a variety of strategies, but their comparison was still at the level of strategy type and not institutional frameworks. Li, Vo, Kowalczyk, Luo and Zhang (2013) and Zhang, Ghenniwa and Shen (2007) improved the adaptive learning mechanisms without trying to encode Harvard, Voss or any identified school.

Research testing abstract concession functions cannot answer the practically useful question, would a procurement manager's outcomes be different if he or she learned the SMARTnership model instead of the Harvard model? This study is an attempt at a direct encoding, but it is also recognised that there are interpretative choices to be made in translating a behavioural philosophy into computational parameters, which limits the claims that can be made.

In the field of pharmaceutical price negotiation, Syversen et al. (2024) find some overlap, as they find that procedural transparency was consistently linked to higher rates of agreement and lower rates of post-agreement disputes across countries. That result in a human-participant setting is broadly consistent with the simulation results reported here, although the mechanisms will almost certainly be different.

2.3 Measuring Negotiation Outcomes

The selection of outcome variables is not innocuous. It embodies assumptions on the purpose of negotiation.

Joint value is an integrative perspective that is the ratio of the available surplus captured by both parties. By this measure, a negotiation that has not resulted in a Pareto improvement is a failure even if both parties are happy with the outcome. It is well grounded in the economics of bargaining (Nash, 1950; Raiffa, 1982) and is suitable for evaluating the efficiency of various bargaining frameworks in exploiting the deal space.

Process quality is not as clear-cut. The measure employed here reflects the symmetry and smoothness of the concession path: the extent to which the sequence of offers is reciprocal and the extent

Table 1. Experimental design summary

to which each party moves. It is a proxy for what practitioners refer to as procedural fairness, which Trötschel, Loschelder, Höhne and Backhaus (2015) showed to have a significant impact on satisfaction of both parties, regardless of the outcome. A negotiation may produce a favourable outcome while still involving irregular or excessive concessions that negatively affect perceptions of the process. The computational measure is only capturing the behavioural pattern and not the subjective experience, and the variable is called concession-process quality to make this distinction. Outcome asymmetry (absolute difference between buyer and seller utility shares) deals with distributional fairness. Two negotiations that result in the same joint value can vastly diverge in allocation: one party might take seventy per cent of the value, or it might be divided evenly. Both outcomes are not necessarily good, but if they are highly asymmetric they tend to breed resentment and renegotiation, especially in repeated interaction situations (Olekals, 1994; De Dreu, Carnevale, Emans, & Van de Vliert, 1995). This feature can help identify frameworks that increase total surplus from frameworks that redistribute it evenly, a distinction that will be shown to have an empirical relevance.

One thing that is not included in this list is a metric of trust as a per-negotiation outcome. The original design was to have a trust index based on concession reciprocity correlations. As explained in section 3.4, that measure was not viable for the purpose it was intended to serve: deterministic concession logic resulted in a constant value in four of five frameworks, reducing the variable to a framework-level label. This is a real constraint of the simulation architecture and is an indication of the need for stochastic trust modelling in future work. Although there is a long history of reputation and trust in the computational literature in multi-period settings (Louta, Roussaki, & Pechlivanos, 2006), the modelling of emergent, experience-based trust in a single ten-round negotiation is still an open design challenge.

3. Methods

3.1 Design Overview

The study was conducted using a within subjects' experimental design. Eight hundred buyer-supplier agent pairs were generated by drawing reservation utilities, aspiration utilities, and issue-level ideal points from uniform distributions, using the Mersenne Twister pseudo-random number generator with two fixed seeds (42501 and 42502), that is, two randomisation blocks of 400 pairs each. There were 4,000 negotiations, each pair negotiating under all five frameworks. Block was added to initial analyses as a control factor and was not significant for any outcome variable (all $p > 0.42$).

All negotiations were conducted in a stylised industrial procurement situation with four interdependent issues: price (weight 0.40), delivery time (0.25), service duration (0.20), and payment terms (0.15). A synthetic construction was used for the scenario, and no proprietary contract data were used.

The design parameters used across all 4,000 negotiations are summarised in Table 1

Parameter	Value
Total negotiations	4,000
Parameter	Value
Frameworks	5 (SMARTnership, Harvard, Conti, Gap Partnership, Voss)
Buyer–supplier pairs	800 (400 per randomisation block)
Design	<u>Within-subjects</u> : each pair negotiated under all five frameworks
Rounds per negotiation	10
<u>Randomisation seeds</u>	42501 (block 1), 42502 (block 2)
Agreement rate	100% across all frameworks
Platform	Python 3.11 (NumPy 1.26, SciPy 1.11)

3.2 Agent Architecture

Every agent has a private utility function defined on the four issues, a reservation utility, an aspiration utility and a concession schedule defined by a framework-specific elasticity coefficient. All frameworks share the same simulation engine, with the only differences being the information-disclosure rule and the different values of the agents' parameters.

Utility function. For agent i and offer vector x , utility is defined as:

$$U_i(x) = \sum_j w_j \cdot v_{ij}(x_j)$$

where w_j denotes the issue weights given in Section 3.1 and v_{ij} normalises issue j linearly onto the interval $[0, 1]$ between the agent's reservation and aspiration points. Joint value is then computed as the mean of the two agents' final utilities, expressed as a percentage of the Pareto-efficient maximum for that scenario instance:

$$\text{Joint Value} = [(U_{\text{buyer}} + U_{\text{seller}}) / 2] \div U_{\text{max}} \times 100$$

Concession schedule. At round t of $T = 10$, the agent's target utility follows:

$$T_i(t) = U_{\text{res}} + (U_{\text{asp}} - U_{\text{res}}) \cdot [1 - (t/T)^{(1/\beta_i)}]$$

where β_i is the concession-elasticity parameter. Higher values of β produce earlier, more linear concessions; lower values produce back-loaded ones. A reciprocity multiplier scales β in response to the opponent's prior concession.

where β_i is the concession-elasticity parameter. The higher the value of β , the earlier and more linear the concessions will be, the lower the value of β , the more back-loaded the concessions will be.

A reciprocity multiplier is a factor that scales β according to the opponent's previous concession.

Offer selection. At every round, the agent looks for the offer in the offer space that maximises the utility of the opponent based on the current belief state and has the target utility. An agent that reveals their weight vector enables the opponent to accurately estimate their weight vector from the disclosure round onwards.

Agreement guarantee. At round 10, all 4000 negotiations reached an agreement. This is not an empirical property of the design. The minimum acceptable share parameter is lower than the minimum possible utility share in all the instances of scenarios, and the concession schedule is designed to guarantee that the target utilities of both agents will converge monotonically to the minimum acceptable ones as t gets closer to T . Thus, an impasse condition cannot occur in the architecture. The simulations would have to be broken down with a walk-away threshold above the convergence point or with an exogenous termination probability, neither of which was done. The uniform agreement rate should thus be interpreted as a limit of the design and not as proof that these frameworks are always able to deliver deals.

3.3 Framework Operationalisation

A parameter configuration and disclosure rule were created for each framework. The translation is an interpretive act, and the results are a description of the performance of these operationalisations and not of all the possible computational renditions of the same philosophy. Table 2 contains the full parameterisation which can be used to replicate or re-specify any agent independently.

Table 2. Computational operationalisation of each framework

Framework	β	Reciprocity	Disclosure round	Min. share	Explore	Claim	Volatility
SMARTnership	1.00	1.40	3	0.35	0.94	0.50	0.040
Harvard	1.00	1.00	2	0.40	0.87	0.50	0.050
Framework	β	Reciprocity	Disclosure round	Min. share	Explore	Claim	Volatility
Conti	0.70	1.00	None	0.40	0.88	0.52	0.045
Gap Partnership	0.80	1.20	3	0.38	0.86	0.53	0.045
Voss (Black Swan)	0.40	0.80	None	0.45	0.78	0.62	0.070

Note. β = concession elasticity; Reciprocity = multiplier applied to β when the opponent concedes; Disclosure round = round at which the agent reveals its issue-weight vector; Min. share = minimum acceptable utility share; Explore = propensity to search alternative offer configurations; Claim = value-claiming propensity; Volatility = stochastic variation in offer generation.

Three design features are worthy of comment. Both the SMARTnership and Harvard agents have the same base concession elasticity ($\beta = 1.00$, which results in linear concession), but with different reciprocity and disclosure: SMARTnership increases concessions to the extent that the opponent does, and it also reveals full issue weights at round 3, whereas Harvard increases concessions symmetrically and it reveals interests but not reservation values at round 2. The Voss agent has a back-loaded concession ($\beta = 0.40$), which means it waits until late rounds to move and is also the agent with the highest minimum-share threshold, indicating tactical anchoring and controlled information flow. No agent was given access to the opponent's concession algorithm or private parameters.

There is a special caution with the operationalisation of tactical empathy in the Voss agent. There is no unambiguous computational analogue of the construct, and here it is modelled as the accurate opponent-preference inference deployed for value claiming. This is one defensible interpretation of the framework, but not the only one, and an alternative specification might change the Voss agent's behavior.

3.4 Outcome Variables

Three dependent variables were computed from the simulation logs at the conclusion of each negotiation.

Joint value captured is the mean of the two agents' final utilities expressed as a percentage of the Pareto-efficient maximum for that scenario instance, as defined in Section 3.2.

Concession-process quality captures the symmetry and smoothness of the concession path on a 0–10 scale:

$$CPQ = 10 \cdot [1 - |\Delta U_{\text{buyer}} - \Delta U_{\text{seller}}| / (\Delta U_{\text{buyer}} + \Delta U_{\text{seller}})] \cdot [1 - \sigma_c / \mu_c]$$

where ΔU denotes each agent's total utility movement across the negotiation, and σ_c and μ_c denote the standard deviation and mean of round-by-round concession magnitudes. The result is floored at zero. Higher values indicate a more reciprocal and less erratic bargaining process.

Outcome asymmetry is the absolute difference between the two agents' final normalised utility shares:

$$\text{Asymmetry} = |U_{\text{buyer}} - U_{\text{seller}}|$$

The index is bounded [0, 1], with zero denoting perfectly symmetric division. Lower values indicate a more balanced allocation of value.

In the simulation logs, a fourth variable, the trust index, was calculated as $10 \cdot \max(0, r_c) \cdot (1 - \text{retraction rate})$ with r_c being the correlation between the two agents' round-by-round concession magnitudes. This metric resulted in one constant score from four of the five frameworks: SMARTnership (9.20), Harvard (6.00), Conti (8.20), and Gap Partnership (9.20), with the only within-framework variation being the Voss agent ($M = 7.27$, $SD = 0.43$). This reflects the architecture of the simulation. Deterministic concession schedules produce the same round-by-round concession correlations no matter what endowment draws are made and so the trust formula reduces to a framework-level constant. The

trust index is therefore not an observation-level result, but a design parameter. It cannot be used in any analysis that involves within-group variation, such as analysis of variance or regression, and is only included in Table 3 as a design characteristic. The theoretical implications are discussed in section 5.5.

3.5 Statistical Analysis

Each buyer–supplier pair was tested under all five frameworks, which means that the observations are not independent and between-subjects analysis of variance would be inappropriate. Nonparametric within-subjects tests are used for all primary analyses. Omnibus differences among the five matched groups were examined using the Friedman test. Post hoc contrasts were conducted with pairwise Wilcoxon signed rank tests with the Bonferroni correction for ten comparisons ($\alpha = 0.005$). Effect size estimates were obtained using paired Cohen's *d*, which is the mean difference between the pairs divided by the standard deviation of the differences.

The Levene's test showed that the variances were significantly different among the frameworks on all three variables, justifying the use of nonparametric techniques. The coefficient of variation

Table 3. Descriptive statistics by framework (n = 800 per framework)

Framework	Joint Value (%) max)	Process Quality (0–10)	Asymmetry Index (0 = symmetric)	Trust Index (design-level)
SMARTnership	93.78 ± 3.75	6.33 ± 0.17	0.042 ± 0.031	9.20 (constant)
Harvard	86.94 ± 4.93	9.69 ± 0.23	0.040 ± 0.030	6.00 (constant)
Conti	87.89 ± 4.65	8.36 ± 0.19	0.051 ± 0.037	8.20 (constant)
Gap Partnership	85.85 ± 4.46	7.78 ± 0.20	0.066 ± 0.041	9.20 (constant)
Voss (Black Swan)	78.07 ± 7.16	5.16 ± 0.15	0.097 ± 0.013	7.27 ± 0.43

Note. Values are mean ± standard deviation. The trust index is reported as a design-level characteristic because it shows zero within-framework variance for four of the five frameworks. Its values reflect the deterministic concession-pattern correlation produced by each agent architecture rather than a per-negotiation outcome. See Section 3.4.

SMARTnership had the greatest percentage of available joint value (93.78%) and the lowest coefficient of variation (4.00%), which means that it had the highest mean value and the most consistent

was calculated for each framework. Pairwise mean differences in joint value were calculated and 95% bootstrap confidence intervals (10,000 resamples, percentile method, seed 42) were calculated. The analysis was performed using Python (SciPy 1.11, NumPy 1.26) and verified in R.

3.6 Data and Code Availability

The entire set of raw simulation logs, including offer-by-offer records for 4,000 negotiations, the negotiation-level dataset, the analysis scripts, the simulation configuration file, and a data dictionary are placed in a public repository. A reviewer-access link is provided with this submission, and a permanent digital object identifier will be assigned on acceptance. Random seeds, all framework parameters and command sequence to reproduce all reported statistics are released.

4. Results

4.1 Descriptive Statistics

Table 3 presents the mean and standard deviation for each outcome variable across the five frameworks. The rankings vary across the three dimensions.

performance. Voss captured the least (78.07%) with the highest variability (CV = 9.17%). The ranking changes with respect to concession-process quality: Harvard was the smoothest and most reciprocal concession path (9.69), SMARTnership was fourth (6.33), and Voss last (5.16). In terms of outcome asymmetry, the low values from Harvard and SMARTnership were very similar (0.040 and 0.042, respectively), showing that both frameworks split value fairly evenly, whereas Voss had the most one-sided outcomes (0.097).

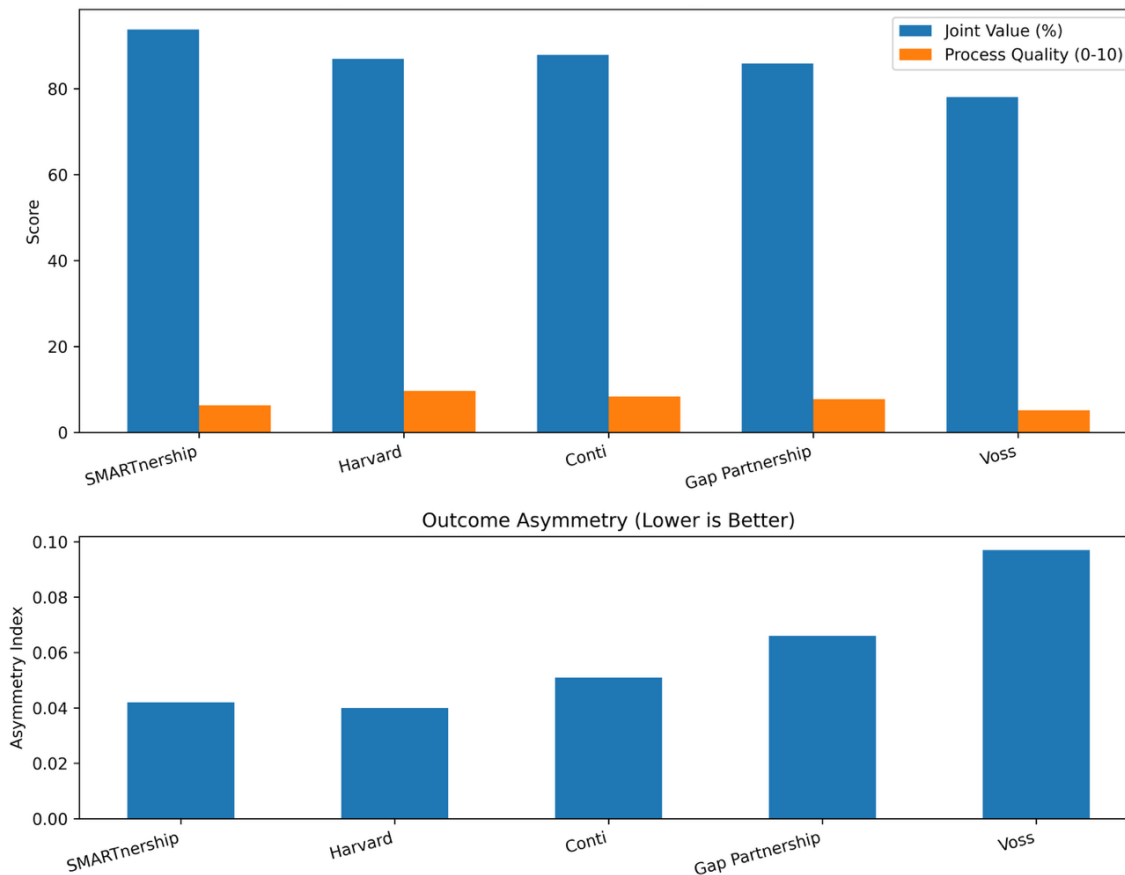


Figure 1. Mean joint value captured by framework, with 95% confidence intervals.

4.2 Omnibus Tests

Friedman tests indicated significant differences across frameworks on all three outcome variables: joint value, $\chi^2(4) = 1,598.04$, $p < 0.001$; concession-process quality, $\chi^2(4) = 3,184.34$, $p < 0.001$; and outcome asymmetry, $\chi^2(4) = 1,086.58$, $p < 0.001$.

Levene's tests confirmed unequal variances across frameworks on all three measures: joint value, $W = 95.50$, $p < 0.001$; process quality, $W = 61.65$, $p < 0.001$; and

asymmetry, $W = 271.58$, $p < 0.001$. These results support the use of nonparametric methods throughout.

4.3 Pairwise Comparisons

Table 4 presents the results of pairwise Wilcoxon signed-rank tests with Bonferroni adjusted p-values and indicates that the SMARTnership improvement over all of the comparison frameworks was statistically significant and practically meaningful.

Table 4. Pairwise Wilcoxon signed-rank tests for joint value (Bonferroni-adjusted)

Comparison	Mean difference	Paired d	Adjusted p
SMARTnership vs Harvard	+6.84	1.09	< 0.001
SMARTnership vs Conti	+5.89	0.99	< 0.001
SMARTnership vs Gap Partnership	+7.93	1.35	< 0.001
SMARTnership vs Voss	+15.71	1.99	< 0.001
Harvard vs Conti	-0.94	-0.14	< 0.001
Harvard vs Voss	+8.87	1.00	< 0.001

Conti vs Gap Partnership	+2.04	0.32	< 0.001
Conti vs Voss	+9.82	1.14	< 0.001
Gap Partnership vs Voss	+7.78	0.90	< 0.001

Note. Mean difference expressed in percentage points of Pareto-efficient maximum. Positive values favour the first-named framework.

After Bonferroni correction all pairwise differences in joint value were still significant. The SMARTnership's advantage over each comparison framework was large (paired $d = 0.99$ to 1.99), and

bootstrap 95% confidence intervals excluded zero in all instances, the comparison to Harvard for example being $[6.40, 7.26]$.

Table 5 presents pairwise comparisons for concession-process quality and outcome asymmetry, where the ranking pattern differs from that observed for joint value.

Table 5. Selected pairwise comparisons for concession-process quality and outcome asymmetry

Comparison	Metric	Mean difference	Paired d	Adjusted p	Direction
Harvard vs SMARTnership	Process quality	+3.36	11.72	< 0.001	Harvard higher
Harvard vs Conti	Process quality	+1.33	4.39	< 0.001	Harvard higher
Harvard vs Gap Partnership	Process quality	+1.91	5.94	< 0.001	Harvard higher
Harvard vs Voss	Process quality	+4.53	16.00	< 0.001	Harvard higher
SMARTnership vs Voss	Process quality	+1.17	5.17	< 0.001	SMARTnership higher
SMARTnership vs Harvard	Asymmetry	+0.002	0.05	1.00	Not significant
Harvard vs Conti	Asymmetry	-0.011	-0.24	< 0.001	Harvard more balanced
Harvard vs Voss	Asymmetry	-0.057	-1.73	< 0.001	Harvard more balanced
SMARTnership vs Voss	Asymmetry	-0.055	-1.63	< 0.001	SMARTnership more balanced

Note. Mean differences for process quality are expressed in points on the 0–10 scale; those for asymmetry are expressed in index units. Negative asymmetry differences indicate that the first-named framework produced more balanced outcomes. The very large paired d values for process quality reflect the small within-framework variance of that metric (see Section 6).

The Harvard agent showed the highest concession process quality

with raw mean differences ranging from 1.33 to 4.53 points on the ten-point scale. On this measure, SMARTnership did better than all but Voss. There was no significant difference between Harvard and SMARTnership on outcome asymmetry (adjusted $p = 1.00$, mean difference 0.002 index units), and both were significantly more balanced than Conti, Gap Partnership, and Voss.

Negotiation Frameworks: Performance Profile Across Three Outcomes

Higher scores indicate better performance on all three dimensions

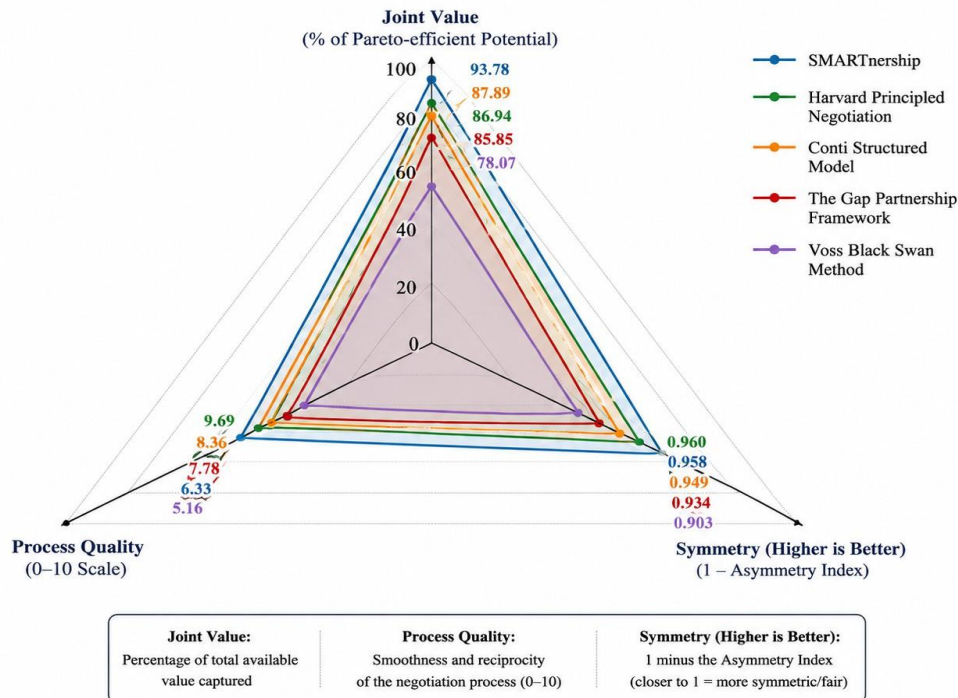


Figure 2. Performance profiles of the five negotiation frameworks across the three evaluation dimensions.

Effect Sizes (Paired Cohen's d) – SMARTnership vs. Other Frameworks

Positive values indicate higher scores for SMARTnership

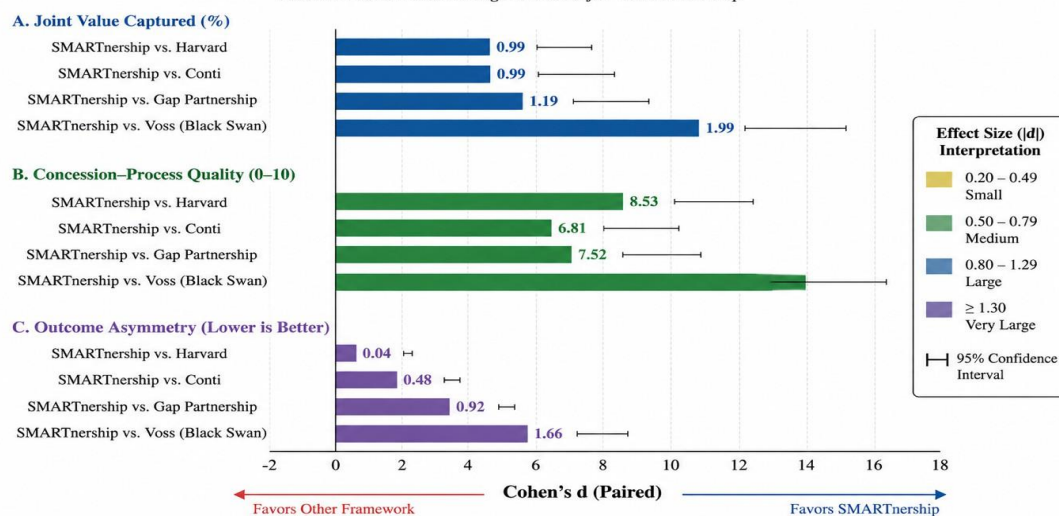


Figure 3. Forest plot of Cohen's d effect sizes for pairwise comparisons involving SMARTnership. three outcome variables.

4.4 Summary of Rankings

Table 6 summarises the framework rankings across the

Table 6. Framework rankings by outcome variable (1 = best)

Framework	Joint Value	Process Quality	Asymmetry
SMARTnership	1st (93.78%)	4th (6.33)	2nd (0.042)
Harvard	3rd (86.94%)	1st (9.69)	1st (0.040)
Conti	2nd (87.89%)	2nd (8.36)	3rd (0.051)
Gap Partnership	4th (85.85%)	3rd (7.78)	4th (0.066)
Voss (Black Swan)	5th (78.07%)	5th (5.16)	5th (0.097)

Note. Harvard and SMARTnership do not differ significantly on outcome asymmetry (adjusted $p = 1.00$); their first and second placements on that dimension are nominal rather than statistically separable.

No framework achieved a top scoring position on all three dimensions. SMARTnership delivered the highest joint value and

was fourth in concession-process quality. Harvard had the most symmetric value division and the smoothest bargaining process and ranked third on value creation. Conti was second in both process quality and value creation and thus had the most consistent overall profile. Gap Partnership finished fourth on two of the three indicators. Voss was last on all the measures.

Framework Ranking Heatmap



Figure 4. Comparative ranking of negotiation frameworks across the three evaluation dimensions.

5. Discussion

5.1 Value Creation and Transparency

The SMARTnership agent has a significant advantage on joint value on all comparisons: 6.84 percentage points higher than Harvard, 5.89 higher than Conti, 7.93 higher than Gap Partnership, and 15.71 higher than Voss with effect sizes ranging from large to very large.

Two features of the SMARTnership agent may explain this result. The SMARTnership agent reveals its complete issue-weight vector at round 3 and finds logrolling opportunities from then on. If the buyer makes it clear that delivery is more important to him than the terms of payment, and the seller's endowment makes delivery cheap to concede, the trade is not available to agents who are concerned only with positional principles. This is the mechanism that Lax and Sebenius (1986) identified as the main source of

integrative gains and the data support this. Once cooperation has been established, the amplified reciprocity multiplier of 1.40 then speeds up mutual concession, shortening the offer sequence, and limiting rounds of positional exchange.

This result should not be interpreted as evidence that transparency will always improve joint outcomes. The benefit is that both agents are subject to the same informational restrictions and that both agents are disclosing the same information. When one party reveals and the other doesn't, as is more likely to happen in the commercial world, the disclosing party could be at a disadvantage. A question that the simulation cannot answer is whether the practical implication is that negotiators should make their priorities known, or that institutions should be developing ways in which both parties make their priorities known.

5.2 Process Quality and the Harvard Agent

The Harvard agent scored the best concession process, beating SMARTnership by 3.36 points and Voss by 4.53 points on the 10-point scale.

The symmetric matching rule in Harvard is that the agent must concede in the same amount as the opponent's last concession. This gives rise to a smooth and reciprocal offer path almost by construction: concessions made by each agent in each round mirror those made by the other agent in the previous round, and the path steadily approaches an agreement. This architecture is well suited to the process-quality metric, which is based on symmetry and low variance of concession magnitude.

The logrolling mechanism of SMARTnership is different. If the agent finds an issue that does not have a high priority for him/her, and a large concession will not cost him/her much, he/she will concede on that issue in one move and ask for a move on an issue of high priority to him/her in return. The resultant offer path is not smooth, but rather a series of step changes on individual issues, which may look like random changes from an outside observer, but may have a cooperative logic. The process-quality metric precisely captures this pattern, which is asymmetric and volatile, even though it is intended to be integrative.

This is a real trade-off which the data cannot settle normatively. A negotiation may be smooth and procedurally fair, but produce less total value, or it may maximise the value but be procedurally unpredictable. In addition, Trötschel et al. (2015) demonstrated that perceived procedural fairness has an impact on negotiator satisfaction beyond the joint-value measure, implying that the Harvard agent's process advantage might be extended to a relational advantage that is not captured by the joint-value measure. Lax and Sebenius (1986) showed that the unchecked smoothness of the procedure systematically results in unclaimed surplus in the absence of preference disclosure. Both propositions may be true at the same time. The point is that the practical implication of a negotiator who cares about the deal and the relationship is that he or she might need a hybrid solution, which the current comparison does not provide.

5.3 The Middle Ground: Conti and Gap Partnership

The two process-oriented frameworks were consistently in the middle for all three outcomes. Conti placed second in both process quality and joint value and had the most stable profile among all frameworks tested. Gap Partnership was ranked third on process quality and fourth on the other two measures.

None of the two frameworks try to maximise joint value by disclosure of preferences, and none try to claim value by psychological leverage. Both opt for predictability – the Conti agent trades off a pre-computed sequence, while the Gap Partnership agent reacts to procedural checkpoints. Moderate performance in all dimensions indicates that the operationalisation of process discipline, as used here, results in adequate rather than outstanding outcomes – adequate for environments where auditability and consistency are important, but not sufficient for the expansion of surplus that disclosure allows or the smoothness of concession that symmetric matching does.

Druckman (2010) suggested that the appropriateness of a negotiation model is determined by situational congruence, and the same rules would work differently in integrative versus distributive, repeated versus single shot negotiations. This study is based on a single scenario and is therefore not able to test that proposition. In a single-issue, high-asymmetry setting where there

is no logrolling benefit from revealing preferences and process discipline is not a concern, it is still possible that the transparency-based models would be outperformed by either Conti or Gap Partnership.

5.4 The Voss Agent and the Limits of Leverage

The Voss agent performed worst on each of the dimensions measured: lowest joint value at 78.07%, lowest concession-process quality at 5.16 and highest outcome asymmetry at 0.097. This was true for all 800 pairs and both randomisation blocks.

The mechanism is simple. Back loaded concessions are created by a concession elasticity of 0.40, which causes the Voss agent to hold off until the late rounds, when the opponent has already conceded most of the times. The next step is to increase the claiming parameter to 0.62, which takes a higher percentage from any surplus. Less total value is created, as the concession pattern of the opponent is not optimally able to react to offers that are late and large, and the value created is split unevenly.

This is in line with the experimental literature on competitive framing. Under competitive and loss framing, negotiators make more deceptive moves and settle for worse joint outcomes than those under neutral and cooperative framing, as Schweitzer et al. (2005) discovered. De Dreu et al. (1995) demonstrated that resistance to concession is greatest in the loss frames, just where cooperation would lead to the most mutual gain. The Voss agent's architecture is an encoded systematic competitive frame which yields the results predicted by the behavioural literature for such a frame. The caveat is that the simulation was bilateral, multi-issue, with equal initial endowments, and a cooperative counterpart was used in four out of five comparisons. Tactical leverage might work better when the opponent is also tactical, in truly distributive situations where the relational costs are not relevant, or in a single instance negotiation. All of these conditions were not tested.

5.5 The Trust Index: A Methodological Finding

The trust index was designed to be a key outcome variable and the empirical test of the idea that trust is a measurable economic input to negotiation performance (Jensen, 2018). That test cannot be performed with simulation data.

Deterministic concession schedules yield the same round-by-round concession correlations for each of four of five frameworks, explaining the constancy of the index. These values do not test anything; the apparent significance is due to the difference between the framework labels, not the per-negotiation outcomes. This is expressly declared because the result was material. The previous analysis resulted in a conclusion that trust accounted for a large portion of variance in joint values, and that this was a by-product of the metric.

The theoretical assumption, that the level of relational trust, if it is indeed different, is a good predictor of the value outcomes, is plausible and has been confirmed in the computational negotiation literature (Louta et al., 2006). Testing it needs stochastic concession logic, to make concession correlations vary within agents, or an experience-based trust model, that would be updated from round to round according to the observed opponent behaviour. Both were not done here. Building such a model and finding out if it will change the framework rankings is the most significant topic on the agenda for future work.

5.6 No Framework Leads on Every Dimension

The most concise overview of the results is provided in Table 6. SMARTnership is the forerunner for joint value. Harvard is first on

concession-process quality, and first on outcome asymmetry, although the difference between Harvard and SMARTnership on that dimension is not statistically significant. Conti has the most well-rounded profile, by dimension. Voss is dead last in all the measures. No framework dominates.

This is more helpful than a pure ranking would have been. It is more plausible on its face that different frameworks will optimise different dimensions, and it gives rise to a feasible selection principle: transparency-based frameworks where the priority is maximising the available surplus, interest-based frameworks where the priority is a smooth and balanced process, and process-oriented frameworks where the priority is procedural consistency and auditability.

What the simulation can't measure, however, is how these frameworks would perform when practiced by seasoned experts under real-world stress, incomplete information, emotional reactions, and reputational consequences. It is those conditions that computational agents eliminate. The rankings reported here refer to the logic of the frameworks operationalised in this context. It is unclear whether they explain the logic of the frameworks as having been experienced in practice.

6. Limitations

Every finding reported here is constrained by the simulation design and merits stating with some precision rather than as a generic disclaimer.

Agents are not human beings, they are rule following systems. They don't worry about giving in first, don't read tone from silence, and don't change strategy if a negotiation is perceived as adversarial. The results present the performance of the encoded logic of each framework under these conditions; to extend the results to trained practitioners, an additional empirical step would be necessary that is not performed in this study.

2. Operationalisation of each framework into a small parameter set requires interpretive judgement and another defensible reading may yield different rankings. The concern is most pronounced for the Voss agent, where tactical empathy, a construct based on emotional attunement and real-time inference, was simulated as opponent-preference inference used for value claiming.

3. The scenario is one stylised procurement negotiation with four issues and fixed weights. No generalization is allowed to distributive single-issue bargaining, multi-party negotiation, cross-cultural negotiation, or service and relationship contracts.

4. 4 frameworks of the trust index were within one of each other, due to the deterministic concession schedules which produce the same concession-correlation pattern, irrespective of endowment variation. It was not possible to analyse trust as an outcome of the negotiation process, and the simulation does not model trust formation as an emergent phenomenon, dependent on the experience of the negotiation process.

5. Quality of concession process is low within framework (CVs between 2.3% and 2.9%), indicating that it is largely a function of agent architecture and not endowment draw. This makes paired effect sizes much larger: the reported d for Harvard versus SMARTnership is actually 11.72, which equates to a raw mean difference of 3.36 points on a 10-point scale; and the d for Harvard versus Voss is 16.00, which equates to a raw mean difference of 4.53 points on a 10-point scale. The ordinal position of this dimension is found to be strong, and the size of the effect

sizes should be interpreted as a reflection of metric determinism rather than a particularly large behavioural difference. For this reason, raw mean differences are reported along with d in Table 5. This is a less severe version of the issue found with the trust index.

6. Concession-process quality is a proxy and not a validated psychometric measure. It does not measure if the parties would have felt the concession process was fair or cooperative, but rather if it was smooth and reciprocal.

At round 10, all negotiations ended with an agreement. The uniform agreement rate is not a finding, but a design property of the architecture, which cannot be in impasse (see Section 3.2). In practice, one dimension of negotiation performance is the ability to avoid breakdown; this is not a dimension on which frameworks can be compared.

8. No sensitivity analysis was performed. The rankings are based on the parameter values in Table 2 and how they would change under reasonable alternative parameterisations is not known. The strength of the ordinal rankings cannot be verified until such analysis is undertaken. The study also lacks external validation with human negotiation data, being the most general limitation on the generality of the results.

7. Conflict of Interest and Author Positionality

The lead author is the developer and chief promoter of the SMARTnership negotiation framework and receives income from training, consulting, and publications related to this framework. SMARTnership's performance on one of the three outcome variables reported in this study is the highest, and is ranked fourth on a second, and not significantly different from the top framework on the third. This represents a substantial conflict of interests and is fully disclosed.

The theoretical framing of the study and the choice of outcome variables are influenced by the author's professional experience negotiating in a collaborative and transparency-oriented manner. One dimension on which a transparency based framework would be expected a priori to perform well is joint value. The study was based on single-issue value claiming, and on variables other than those listed above, such as time to agreement, resistance to manipulation, or negotiator satisfaction, might yield a different ranking.

Three safeguards were used. Each of the five framework modules were coded by separate teams who were not aware of the comparative objectives, and only identified by label. The operationalisations in Table 2 were given to academic colleagues with no commercial relationship to the SMARTnership organisation for review prior to execution, whose names have been provided to the editorial office. All code and raw logs are published, so that any reviewer can re-specify the SMARTnership agent and re-run the comparison. The author welcomes such re-analysis.

There was no outside funding for this research. No sponsor was involved in the design, conduct, analysis or publication of the study.

8. AI Disclosure Statement

The negotiation agents are rule-based multi-agent software written in Python and are a method of the study, not an authoring aid; no generative language model produced negotiation behaviour or outcomes. Generative AI coding assistant was used to generate and refactor parts of the simulation and analysis scripts which were reviewed, tested, and verified before deposit; the verified version

is the one included in the released repository. No generative tool was used to select tests, compute results or interpret output; statistical analysis was carried out in Python and verified in R. Language editing and readability improvement were carried out using a generative AI tool during manuscript preparation and all the language edited/revised text was checked and revised by the author. No AI tool is cited as an author, and the author accepts full responsibility for the integrity, accuracy and originality of the manuscript.

9. Future Research

These limitations suggest several directions for future research. The next most important step is external validation. Testing practitioners in each framework using a similar multi-issue procurement negotiation, using experimental controls and sufficient sample size, would determine whether the computational rankings are valid once human skill and emotion, along with real-time inference, are introduced into the process. The matched-pair design employed here might be realized in a human study by showing the same instances of a scenario in the same framework conditions and the results of the simulation can serve as a concrete comparison.

The need for a stochastic trust model is closely related. The trust index is not just a data glitch; it's a missing piece in the simulation's design. The negotiation trust that is established is a gradual process that is based on observation, inference, and experience, and differs significantly from one dyad to another, even if the overall strategy is the same. To test the study's central theoretical claim, deterministic concession schedules could be replaced with learning-based or stochastic ones, and an experience-dependent trust metric could be introduced that would vary over the course of the rounds.

Sensitivity analysis is a more immediate need. A systematic variation of each of the framework parameters, with the reporting of whether and to what extent the ordinal rankings change, would set up the extent to which the current findings are robust to the interpretive choices made in Table

2. This analysis could be done with the current code base and the released logs and would greatly enhance the evidentiary value of the results.

The one-scenario design can be replicated by others in multiple scenarios. Distributive negotiations, in which interests are opposed rather than across issues, may reward different framework logics entirely; tactical anchoring may be more effective than collaborative disclosure where there is no logrolling. A bilateral procurement simulation doesn't account for the dynamics of multi-party negotiations, service contracts, or cross-cultural dyads.

Independent re-implementation of the framework agents by scholars not related to any of the five frameworks would be most direct to the operationalisation concern. The more researchers without vested interests encode the frameworks from their published descriptions and the more agents are generated that rank them the same, the more confidence is gained in the results. Where they do give different rankings, the sensitivity of the results to operationalisation choices is quantified rather than just acknowledged.

There is a longer-term potential with NLP and large language model agents. Some frameworks make the linguistic aspects of negotiation a key part of the process, but rule-based agents are incapable of doing so. An agent capable of creating calibrated

questions and tactical empathy statements directly, without approximation via a claiming parameter, would allow for a truer test of the Black Swan framework and may yield a very different result.

10. Conclusion

Negotiation frameworks are advocated based on the practitioner testimony and the theoretical argument. Comparative empirical evidence with everything else equal but the framework itself has been almost completely lacking. This study provides such evidence for five popular frameworks, based on 4,000 matched-pair agent simulations of the same 800 buyer-supplier scenarios negotiated under each set of rules.

The results don't yield a sole winner, and that's the most significant. SMARTnership was able to capture much more joint value than any comparison framework, with large to very large effect sizes for all four contrasts. The mechanism is traceable: early preference disclosure helps both agents to discover logrolling trades that cannot be reached by positional bargaining, and an amplified reciprocity response facilitates agreement once cooperation is reached. The data suggest a transparency-based approach for a negotiator whose main goal is to increase the total surplus.

Harvard Principled Negotiation resulted in the most even and mutually beneficial concession and the most nearly even split of value. It has an asymmetric matching rule which produces a mechanically fair round-by-round bargaining path. On outcome asymmetry, both Harvard and SMARTnership split the value in half, but they got there in different ways and at different levels of total value.

The Voss agent was worst on all dimensions measured. Back-loaded concessions and high claiming parameters take a bigger bite out of a smaller surplus, and the simulation showed no circumstances where such a trade-off was net positive. This does not mean that the Black Swan Method is not of value in practice, as tactical leverage might work differently in single-issue, time-constrained, or truly oppositional situations not considered in the design. It does show that in a multi-issue commercial negotiation with a fully matched counterpart, the way it was operationalised here did worse on all measures used.

Two results deserve to be noticed outside of the negotiation literature. The first is a methodological one. As a central variable and as an integral part of the original theoretical approach, the trust index remained stable within four out of five frameworks. This wasn't a measurement problem, but the simulation showing that deterministic concession logic does not produce emergent trust dynamics. To claim that trust is a within-framework per-negotiation outcome, it is necessary to first show within-framework variance in that outcome.

The second is substantial. The value-creation-versus-concession-process-quality trade-off is a genuine one and is in the opposite direction to that of the original framing of the study. The framework maximising value came out fourth on process quality and the framework maximising process quality came out third on value. These are not artefacts. They are a true compromise of what integrative frames optimise: logrolling moves that make the surplus bigger seem counterintuitive; smooth symmetric concessions that seem cooperative leave potential trades unexploited. The practical contribution of this research is in resolving that tension, using hybrid methodology or selecting a context-specific framework.

There is no simulation that really addresses the negotiation process

between people. Agents follow rules; practitioners follow instincts, based on culture, relationship history, risk tolerance and pressures of the specific deal. The difference between the modelled system and the human system is real and recognized. What the simulation offers is a statement of what each framework's logic yields when it is applied to scale, without error, under the same conditions. It is a question this study is intended to inspire rather than answer: Can practitioners keep up with that baseline, and will the rankings hold up with the addition of human factors?

All data and code are made available. The invitation to re-specify the agents, change the parameters, and test to see if the rankings hold up is intentional.

References

1. Amgoud, L., Dimopoulos, Y., & Moraitis, P. (2007). An abstract framework for argumentation-based negotiation. In *Proceedings of the 20th International Joint Conference on Artificial Intelligence* (pp. 1456–1461). IJCAI.
2. Conti, G. (2019). *The psychology of negotiation*. Geneva Business School.
3. De Dreu, C. K. W., Carnevale, P. J., Emans, B., & Van de Vliert, E. (1995). Outcome frames in bilateral negotiation: Resistance to concession making and frame adoption. *European Review of Social Psychology*, 5(1), 97–124. <https://doi.org/10.1080/14792779443000021>
4. Debenham, J. (2003). An eNegotiation framework. In M. Bramer, A. Preece, & F. Coenen (Eds.), *Research and development in intelligent systems XIX* (pp. 79–92). Springer. https://doi.org/10.1007/978-1-4471-0643-2_6
5. Debenham, J. (2004). Multi-issue bargaining in an information-rich context. *Knowledge-Based Systems*, 17(2–4), 97–105. <https://doi.org/10.1016/j.knosys.2004.03.010>
6. Druckman, D. (2010). Frameworks, cases, and experiments: Bridging theory with practice.
7. *International Negotiation*, 15(1), 1–23. <https://doi.org/10.1163/157180610X506947>
8. Fisher, R., & Ury, W. (1981). *Getting to yes: Negotiating agreement without giving in*. Houghton Mifflin.
9. Friedman, M. (1937). The use of ranks to avoid the assumption of normality implicit in the analysis of variance. *Journal of the American Statistical Association*, 32(200), 675–701. <https://doi.org/10.1080/01621459.1937.10503522>
10. The Gap Partnership. (2023). *Advanced negotiation programme manual*. The Gap Partnership.
11. Hausken, K. (1997). Game-theoretic and behavioral negotiation theory. *Group Decision and Negotiation*, 6(6), 511–532. <https://doi.org/10.1023/A:1008684225781>
12. Jensen, K. (2018). *Negotiation: The SMARTnership approach to value creation*. Wiley.
13. Lai, G., & Sycara, K. (2009). A generic framework for automated multi-attribute negotiation. *Group Decision and Negotiation*, 18(2), 169–187. <https://doi.org/10.1007/s10726-008-9119-9>
14. Lax, D. A., & Sebenius, J. K. (1986). *The manager as negotiator: Bargaining for cooperation and competitive gain*. Free Press.
15. Li, M., Vo, Q. B., Kowalczyk, R., Luo, X., & Zhang, M. (2013). Automated negotiation in open and distributed environments. *Expert Systems with Applications*, 40(15), 5931–5942. <https://doi.org/10.1016/j.eswa.2013.05.033>
16. Louta, M., Roussaki, I., & Pechlivanos, L. (2006). Reputation-based intelligent agent negotiation frameworks in the e-marketplace. In *Proceedings of the International Conference on e-Business* (pp. 5–12). <https://doi.org/10.5220/0001426600050012>
17. Nash, J. F. (1950). The bargaining problem. *Econometrica*, 18(2), 155–162. <https://doi.org/10.2307/1907266>
18. Olekalns, M. (1994). Context, issues, and frame as determinants of negotiated outcomes. *British Journal of Social Psychology*, 33(2), 219–232. <https://doi.org/10.1111/j.2044-8309.1994.tb01018.x>
19. Raiffa, H. (1982). *The art and science of negotiation*. Harvard University Press.
20. Schweitzer, M. E., DeChurch, L. A., & Gibson, D. E. (2005). Conflict frames and the use of deception: Are competitive negotiators less ethical? *Journal of Applied Social Psychology*, 35(10), 2123–2149. <https://doi.org/10.1111/j.1559-1816.2005.tb02212.x>
21. Sierra, C., Jennings, N. R., Noriega, P., & Parsons, S. (1997). A framework for argumentation-based negotiation. In M. P. Singh, A. Rao, & M. J. Wooldridge (Eds.), *Intelligent agents IV (LNCS Vol. 1365)*, pp. 177–192. Springer. <https://doi.org/10.1007/BFb0026758>
22. Spector, B. I. (1977). Negotiation as a psychological process. *Journal of Conflict Resolution*, 21(4), 607–618. <https://doi.org/10.1177/002200277702100404>
23. Syversen, I. D., Schulman, K., Kesselheim, A. S., & Zhang, C. (2024). A comparative analysis of international drug price negotiation frameworks: An interview study of key stakeholders. *The Milbank Quarterly*. Advance online publication. <https://doi.org/10.1111/1468-0009.12714>
24. Trötschel, R., Loschelder, D. D., Höhne, B., & Backhaus, K. (2015). Procedural frames in negotiations: How offering my resources versus requesting yours impacts perception, behavior, and outcomes. *Journal of Personality and Social Psychology*, 108(3), 365–384. <https://doi.org/10.1037/pspi0000009>
25. Voss, C., & Raz, T. (2016). *Never split the difference: Negotiating as if your life depended on it*.
26. HarperBusiness.
27. Zhang, Z., Ghenniwa, H., & Shen, W. (2007). A framework for adaptive negotiation in multi-agent systems. In *Proceedings of the 11th International Conference on Computer Supported Cooperative Work in Design* (pp. 1043–1048). IEEE. <https://doi.org/10.1109/CSCWD.2007.4281470>
28. Zhang, Z., Shen, W., & Ghenniwa, H. (2008). An adaptive agent negotiation framework. In W. Shen, J. Yong, Y. Yang, J. Barthès, & J. Luo (Eds.), *Computer supported cooperative work in design IV (LNCS Vol. 5236)*, pp. 216–226. Springer. https://doi.org/10.1007/978-3-540-92719-8_22