

From 2D Object Space to 3D Perspective Structure

Overview

This paper addresses the structural conditions under which objectivity can emerge when multiple perspectives are present. Standard approaches typically presuppose a shared object space and treat differences between observers as secondary deviations. This working paper suspends that assumption and formalizes the transition from a two-dimensional object space to a three-dimensional perspective structure.

The formalization introduces a perspective group that acts on objects, allowing perspective-invariant comparison through alignment operations. This mathematical framework generates operators for reflexivity and contextual coherence without requiring external assumptions. The extension from D2 (systematic comparison) to D3 (perspectival alignment) and D4 (contextual integration) provides a dimensional logic that stabilizes cooperation in rational decision-making.

This compact technical paper demonstrates the core mechanism of dimensional logic through a specific application: perspective alignment in cooperation problems. It serves as the formal foundation for the framework developed in "Adaptive Rationality through Dimensional Logic" (Stegemann, 2025) and illustrates how perspective-invariant metrics can be operationalized in decision architectures.

1) Basic Idea: From 2D Object Space to 3D Perspective Structure

D2 (Object Space): A cognitive state space $X \subseteq \mathbb{R}^2$ (e.g., two salient criteria of a problem).

Perspective Space P: Possible "viewpoints" on X . Formally: a group or manifold of transformations G (e.g., rotations, scalings, re-weightings) that act on X .

D3 (Total Space): The perspective fiber bundle $E = X \times P$ (with trivial bundle) with projection $\pi : E \rightarrow X$. A point in E is a pair (x, p) : "object x from viewpoint p ".

Thus D3 is not a new rule, but rather an extension of the space: We attach to each object x its possible viewpoints p .

2) Perspective-Invariant Comparisons (the Bird's Eye View formally)

Errors in D2 often arise because two actors see the same object x in different frames. In D3 we do not compare raw x_1, x_2 , but rather aligned:

$$d_G(x_1, x_2) = \inf_{g \in G} d(g \cdot x_1, x_2)$$

where G is the perspective group (e.g., $SO(2)$ for rotations, or diagonally weighted scalings), d is a metric on X (e.g., Euclidean), and d_G is the perspective-invariant distance: We first allow an adjustment of the view (alignment), then measure the proximity.

Interpretation: D3 implements the "bird's eye view" by minimizing over all viewpoints. The result is less error of mismatch and thus fewer coordination errors.

Technically cleaner: For compact G (e.g., $SO(2)$) the infimum exists; it is continuous in x_1, x_2 ; often a minimizer g^* exists.

3) From Perspective to Decision: Lift-Align-Project

A generic decision procedure on D3:

First, **Lift**: Raise the case into the fiber bundle: $(x, \cdot) \in X \times P$.

Second, **Align**: Find the best perspective $g^{**} = \arg \min_{g \in G} d(g \cdot x_I, x_{You})$ or relative to a norm/rule x_{Norm} .

Third, **Project**: Decide in the aligned frame via invariants or smaller distance.

This pipeline is purely mathematical (no ad-hoc operators needed). "Reflexivity" appears as a special case when x_{You} is one's own model of the other; "context coherence" appears when x_{Norm} is a contextual reference object. Both are now derived, not postulated.

4) Why Cooperation Works (without Operators)

Many defection motives are frame-mismatch costs. In D2 they act as disturbance; in D3 the alignment reduces these costs. Formally you can reduce the earlier threshold

$$T - R \text{ ("advantage of defection")}$$

by an alignment gain $A_G \geq 0$:

$$\text{Cooperation rational} \Leftrightarrow R + A_G(x_1, x_2) \geq T$$

With

$$A_G(x_1, x_2) \equiv \gamma \cdot (d(x_1, x_2) - d_G(x_1, x_2))$$

that is, "how much distance disappears through perspective alignment". $\gamma > 0$ scales the significance of this gain.

Clou: A_G is not an exogenous operator term; it emerges from the geometry.

5) Concrete, Minimal-Operational Example

D2-states: $x = (u, v)$ (e.g., short-term payoff vs. reputation).

Perspectives G : Rotations $R_\phi \in SO(2)$ ("which axis is more important?").

Alignment: Find $\phi^{**} = \arg \min_\phi \|R_\phi x_A - x_B\|$.

Decision: Cooperate when

$$R + \gamma \cdot (\|x_A - x_B\| - \min_\phi \|R_\phi x_A - x_B\|) \geq T$$

Interpretation: If a pure frame alignment already eliminates large discrepancies, cooperation becomes rational without ever defining $\Delta\mu/\Delta\kappa$.

6) Connection to Your Previous Operator Model

Operators $(\Delta\mu, \Delta\kappa)$ are now projections of the D3-geometry into D2:

$\Delta\mu \approx$ alignment gain between agent frames (self/other), arising from $\min_{\{g \in G\}}$.

$\Delta\kappa \approx$ alignment to a context reference frame (norm/institution), likewise from $\min_{\{g \in G\}}$.

Advantage: You do not need to postulate them; they are computed byproducts of the perspective extension.

7) How to Actually Implement This (Simulation/LLM)

In PD simulation:

Replace the raw run-to-decision with proximity of goals through d_G (or use A_G as additional utility). This is a short alignment step (with $SO(2)$ even closed-form solvable).

In LLM scenarios:

"Perspectives" are frame parameters in the response space (weighting of intention vs. policy, style, scaling of risk vs. reputation). Alignment means: transform answer candidate so that it fits user frame or norm frame (e.g., through re-ranking via d_G or search-based prompt refactoring).

8) Next Steps (concrete & start small)

First, **choose G** : initially $SO(2)$ (only rotations), later $GL(2)^+$ (rotation + scaling).

Second, **define d** : Euclidean or Mahalanobis.

Third, **implement invariant d_G** : program (simple scan over ϕ or closed solution).

Fourth, **replace** in existing code the previous "reflexivity term" with A_G .

Fifth, **test**: Heatmaps & learning curves with/without perspective alignment compare.