

How "Ought" Derives from Evaluatively Structured "Is" in Cognition

Bree Beal

2026 - Formal Specification 1.6.3

Revision 1.6.3: makes three bounded formal-integrity corrections without reopening the theory architecture. First, the exact subject-binding node that converges with the relationship and object-binding nodes must itself locally actually contribute to the token-relative evaluation; a second subject node cannot supply that contribution. Second, a strict-bender classification now requires the represented relationship profile to continue actually contributing in the registered path and removal counterpart states while leaving full output and moral-classification invariance unrequired. Third, the current dependency audit enumerates every governing-head argument and fails closed on an unknown or untyped variable instead of silently ignoring it. These corrections preserve the three-builder architecture and the nonexhaustive `CoreChangingFactor_s` rule. This is an unratified, unmerged RC4 candidate.

Revision 1.6.2: makes two bounded artifact-integrity corrections without changing the theory architecture, formal predicates, or registered claim status. First, it removes a residual ambiguity in the empirical discussion of strict bending: a strict bender may alter an action description or another registered downstream locus while preserving the exact core, but a change to an agent, object, or other core-individuating construal is core-changing. Second, it limits the relation-to-existing-work chapter to classifications governed by the current comparison Dataset. Pollock, Horty, Mikhail, Cushman/Crockett, Gibbard, and deontic-argumentation comparisons remain explicitly deferred candidates rather than active comparison conclusions. The governed Dung grounded-semantics comparison, the hierarchical Bayesian moral-weight comparison, and the bounded aggregate moral-psychology conclusion are retained. This is an unratified, unmerged RC3 candidate.

Revision 1.6.1: repairs comparison-class typing without changing the intended architecture. One explicit comparison class is threaded through the valence, core, relationship-strength, builder-qualified, and moral-relationship predicates. Token-relative relationship contribution and strict-bender conditions constrain it to `ComparisonClassOf(kappa,c)`, and `MoralJudgment` binds that same prospective class before applying either relationship test. The revision also corrects the strict-bender example so that action description, not a core-individuating construal, is the downstream locus, and standardizes the reader-facing builder gloss as perceived agency. This is an unratified, unmerged RC2 candidate.

Revision 1.6.0: records the Gate 2 moral-relationship reorganization. It replaces the single positive final-value coordinate with one inherent-valence profile whose positive and negative components remain independently represented; the profile is never reduced to a net scalar and the positive-only final-value boundary preserves the corresponding 1.5.3 branch verdict. A derived relationship-strength profile then preserves separate positive-by-agency and negative-by-agency interaction branches; its positive-only slice preserves the corresponding 1.5.3 relationship-strength verdict without reinstating a one-scalar aggregate or a fourth builder. It replaces Closeness with valence-neutral, non-scalar Relational Positioning and closes the governed route taxonomy over identity-mediated positioning, intimacy, and direct encounter. Identity-mediated positioning has shared, contrastive, and complementary-role forms; an institutional or otherwise impersonal directed assignment qualifies only through an independently specified, operative identity or role schema that locates the exact object relative to the exact subject and contributes through that bridge. Functional access remains separate. Represented distance is absence of Relational Positioning under an adequate detection scheme, not an additional position; opposition is a contrastive-identity form and does not entail negative inherent value.

Revision 1.6.0 also narrows *bender* to a core-preserving downstream operation on a live, actually contributing path. Its removal intervention must preserve an exact before/after snapshot of the bound core. A factor that changes a relationship builder is typed instead as a builder realizer, builder input, or builder modulator for that operation. One representation may receive separately evidenced valence, positioning, strict-bender, and attractor edges; no label supplies another edge by association. Vulnerability, power, and sentience from the 2021 Moral Space proposal remain candidate bender sources only. This revision retains the Gate 1 contribution, system-binding, typed-integration, uncertainty, praise, satisfaction, and override repairs.

Revision 1.5.3: performs the ratified Gate 1 formal-integrity repair. It types uncertain-prospect ordering, removes universal strong semantic CET from NDT's assumptions, separates encoded-evaluativity existence from non-negligible occurrence, and guards output and moral predicates by an evaluable, uniquely system-bound model instance. It also distinguishes structural override licensing from operative override, causal activity from surviving support and actual contribution, and relationship-level profile NESS from the independently required local-contribution verdict. The revision prospectively indexes contribution schemes to declared comparison classes, repairs praise and satisfaction definedness, qualifies relationship-witness deletion, and adds the narrow same-chain integration condition supported by the Gate 1 countermodel audit. It changes no relationship builder and introduces no new moral-positioning subtype.

Revision 1.5.2: binds builder-profile contribution to a controlled, same-chain witness and excludes coincidental covariation through unrelated relationship features. It retains redundant support through a profile-level minimal-sufficient-set (NESS-style) clause rather than requiring every builder or Close realization to be an individual but-for cause. It also sharpens Closeness as a non-scalar family of independently specified subject-relative relational roles, governs the registration of new

routes, and indexes perceived agency to a declared modal-temporal horizon while separating response affordance from present feasibility. No change makes functional access a builder or collapses the three-builder architecture.

Revision 1.5.1: closes the token-relative evaluation and polyphony definitions over their displayed variables; repairs the sufficient bridge for conditional cross-scheme convergence by defining and preserving the difference-maker verdict over a nonempty intervention basis; and makes functional access and bound builder-profile contribution explicit in the local-contribution chain. The exact agent, object, referent, time, and frame tuple is threaded through that contribution test, and the typed vocabulary distinguishes support units, processes, interventions, schemes, maps, and route witnesses. It also clarifies the relationship architecture after the 1.5.0 decomposition: Close remains the third overall builder, but it is a non-scalar, multiply realized role rather than a quantitative coordinate. Shared identity, intimacy, and encounter/empathic uptake are named, nonexclusive, individually optional, and nonexhaustive Close realizations; standing-mediated and other realizations remain possible. Moral type requires some Close realization, not any particular named subtype. Functional access remains enabling contribution machinery rather than a Close realization.

Revision 1.5.0: retired existential closeness as a universal scalar quantity. Final value and perceived agency remained the two noncompensable quantitative relationship roles. Interface fit and availability to memory, attention, or awareness move into functional access and actual-contribution machinery. Shared identity, intimacy, and encounter or empathy were separated as optional, nonexclusive, nonexhaustive provenance-route forms, alongside mixed and unclassified forms. Revision 1.5.1 retains that decomposition but clarifies that the umbrella Close role, multiply realized by such forms, remains a necessary third builder. The moral token still requires a bound qualifying relationship, the subject and object gates, and profile-sensitive local actual contribution.

Revision 1.4.9: restores explicit empirical answerability without treating disputed lexical verdicts as self-interpreting falsifiers. It states a direct represented-subject pressure test, restores an explicit commitment to restrict, revise, or abandon an explication after sustained failure on independently specified tests, and registers full-architecture instantiation and provenance-based kind unity as empirical bridge hypotheses. It also recasts the builder argument as conditional and contrastive, distinguishes four evaluative levels, and treats practical tractability as a cross-cutting constraint. The title now names the evaluatively structured input explicitly; this is a clarification of the existing is-to-ought claim, not a retreat from it. No formal predicate, executable behavior, or Stage 6 status changes.

Revision 1.4.8: clarifies the three relationship builders as constitutive role-conditions of NDT's proposed explanatory kind while leaving their indicators, decomposition, quantitative form, population distribution, and possible extension open. It distinguishes an explication-adequacy challenge from ordinary hypothesis falsification, reframes empirical work as an open development program rather than a global criterion of theoretical value, and preserves without qualification the claim that NDT describes how ought is inferred from evaluatively structured is in cognition. No formal predicate, is-ought clause, executable behavior, or Stage 6 status changes.

Revision 1.4.7: synchronizes the bounded character-appraisal rival with the full prospective/reconstructive grammar. The rival may apply to some third-person and first-person reconstructive cases; first-person prospective judgment remains the discriminating control. It also registers Executable Partial Reference Model 1.6 as the finite companion that implements the 1.4.6 local-contribution repair. No constitutive predicate or Stage 6 status changes.

Revision 1.4.6: defines local evaluative contribution for OrganizesOn and ObjectContributes_s under the same fixed constitutive scheme used by token-level actual contribution; mere reachability is no longer sufficient. It also registers kind-adequacy criteria for the explanatory construct moral cognition, distinguishes three empirical levels, limits a character-appraisal alternative interpretation, and consolidates one restrained source-sensitivity lineage note. No other 1.4.5 behavior is changed.

Revision 1.4.5: repairs two specification-level gaps detected by the Phase 1 derivation extractor. Section 6 used ReconstructivelyAdmissible_s and SurvivesOn without defining either symbol. Section 6.1 now defines reconstructive admissibility by tying a token's base decision model and frame to ReconstructionFormed_s; Section 6.2 now defines route-specific survival as LiveRoute on the specified alignment route. No existing expression changes behavior, and all other 1.4.4 commitments remain adopted.

Revision 1.4.4: makes four narrow authority-level repairs. It replaces metaphysical-sounding “uncaused” examples with representationally agentless cases; states the logical independence and conjunctive necessity of the component minimums and combined relationship threshold; requires ControlRecruited to be identified through a pre-specified control-state indicator not defined by the target output; and updates the empirical-status notes and registry crosswalk. All other 1.4.3 commitments, including relation qualifiers, ownership, and conditional cross-scheme convergence, remain adopted.

Abstract

Normative Detachment Theory (NDT), also called the Is-to-Ought model of normative cognition, is a descriptive architecture of how cognitive systems form prescriptions and appraisals. It distinguishes semantic success, represented

outcome evaluation, and frame-relative practical guidance. In the encoded case, semantic statuses connect to outcome distributions through an independently identified control profile and a pre-specified aggregation rule rather than through a hidden desire or a theorist-imposed welfare ranking; a decision model then maps feasible actions to represented outcomes, standards, priorities, and a violation-sensitive fallback structure.

The primary output unit is a judgment-derivation token with an event anchor and an actual functional support graph. A candidate becomes an output only when a provenance-closed frame-model-token chain is live, anchored, and actually contributing within that graph. Actual contribution is evaluated inside a fixed constitutive model instance whose realization-profile scheme preserves independently specified functional organization over a declared comparison class. Methodological admissibility does not make a scheme true: invalid instances and unresolved rival schemes with divergent verdicts withhold application rather than forcing a negative token verdict. Constitutive adequacy alone does not guarantee cross-scheme convergence; contribution equivalence follows conditionally where adequate schemes are connected by a same-resolution, sufficiency-preserving bridge. Redundant routes may both contribute; a live route that changes no constitutively relevant realization feature remains profile-silent and cannot confer moral type.

Moral cognition is a grounded subset of normative cognition. Every moral token represents at least one moral subject and at least one moral object, although the subject may be implicit in the surface sentence. A token is moral only when a frame-matched represented moral relationship, a route-carried subject binding, and the relevant object binding jointly reach the token-relative evaluative structure along one actually contributing live chain. The relationship combines an inherent-valence profile, perceived agency, and Relational Positioning. Positive and negative inherent value are independent components rather than opposite ends of one scale. Relational Positioning is valence-neutral and is realized by an identity-mediated, intimacy, or direct-encounter route. Appraisals of events represented as having no agent behind them - for example, representing an earthquake as terrible while attributing it to no subject - may be axiological, tragic, or otherwise normative, but they are not moral judgments. Institutions and collectives qualify only when the system represents them, or their members, as agents.

The practical grammar includes prospective guidance, violation, wrongness, blame, praise, and generic positive or negative appraisal. Surface-agentless statements such as "what happened was wrong" receive moral type only when their actual derivation contains an explicit or implicit moral subject linked to the appraised target. Frame and judgment-token identity are explicit. Polyphony requires distinct frame activations, a common case, local coherence, and the comparison bridge appropriate to the pair. A shared execution map is not sufficient when appraisal dimensions still lack a shared question or surviving dimension map. Such cases remain unaligned plurality rather than conflict.

Comparative explication adequacy, formal countermodel structure, operational identification, and calibrated measurement are kept distinct; practical experimental tractability is a cross-cutting constraint rather than a fifth evidential level. Provenance is a latent functional relation; self-report is not privileged, and tests must triangulate interventions, transfer, process measures, and model comparison. The accompanying executable partial reference model checks finite structural consequences and boundary cases; it does not constitute independent formal verification, empirical confirmation, or a mechanized proof of equivalence between prose and code.

Contents

- 0. Scope and Commitments
 - 1. The Set-Theoretic Hierarchy
 - 2. The Constitutive Evaluativity Thesis
 - 3. The Normative Detachment Thesis
 - 4. The Moral Detachment Thesis
 - 4a. The Perceived Moral Relationship
 - 4b. From Moral Relationship to Moral Judgment
 - 4c. Worked Composition Tests

0. Scope and Commitments

Because the theory is easy to misread as more ambitious, more metaphysically committed, or more behaviorally deterministic than it is, the following commitments fix the frame for everything that follows.

0.1 Level of analysis

NDT states a computational architecture together with abstract functional-representational constraints on any implementation of that architecture. It specifies the transformation achieved - from framed representation and evaluation to indexed judgment - and the dependency structure that must be preserved for provenance, defeat, and moral typing to be well defined. Routes, premises, warrants, imports, and support graphs are therefore functional roles, not a commitment that the brain or another system literally stores labeled graph objects. Any neural, symbolic, dynamical, affective, perceptual, or hybrid implementation that preserves the relevant dependencies counts as an implementation. NDT is consequently more committal than a purely extensional computational theory but less committal than a specific algorithm or neural model.

0.2 Relation to Hume

NDT does not refute Hume's Guillotine. If "is" is stipulated to contain no evaluative ordering or standard of any kind, no prescriptive conclusion follows from it merely by description. The claim is instead that cognitively instantiated descriptions can fail to satisfy that stipulation. They may possess success conditions and may be connected by standing architecture to represented outcomes that matter to the system. The theory concerns the actual contents and transformations of cognition, not a deduction from evaluatively empty propositions.

0.3 Relation to belief-desire psychology

Where the evaluative input is an explicitly tokened goal, desire, command, or policy, NDT restates the practical syllogism and is compatible with ordinary Humean psychology. The distinctive case is encoded: an outcome ordering is available without a separately tokened conative state because a stake-sensitive consumer role is constitutive of the relevant representation. The theory does not infer an action directly from representational correctness. Section 2 separates semantic ordering, outcome ordering, and practical closure; Section 3 supplies the frame-relative decision model that ranks actions.

0.4 Descriptive commitment

The theory is descriptive at every level. It classifies the structure and provenance of judgments without deciding whether their contents are objectively correct. Descriptive neutrality concerns the correctness of represented verdicts; it does not prevent NDT from advancing substantive hypotheses about the psychological structure by which those verdicts acquire normative or moral type. A token may be moral in the psychological sense defined here while being rejected by the theorist, another community, or another explicitly specified normative framework. Conversely, a token may reproduce the language or behavior associated with Kantian, utilitarian, religious, institutional, or political morality while remaining nonmoral if no represented moral relationship organizes its actual evaluative structure.

This consequence is intentional. NDT does not grant moral status to a philosophical label, a professed principle, or an outwardly impartial style of reasoning. Some self-described consequentialist or deontological reasoning may therefore be normative but nonmoral. Abstract and collective objects such as humanity, rational agency, sentient welfare, a polity, or a future generation can ground moral relationships, but only when the relevant relationship token, builder profile, object binding, and provenance are independently instantiated rather than supplied post hoc to save a preferred classification.

0.5 Constitutive and implementation claims

Constitutive claims state what a modeled phenomenon consists in and are load-bearing within the theory. Implementation claims specify one possible computational realization and are replaceable if a better realization preserves the constitutive structure. The three-builder relationship condition - an independently two-component inherent-valence profile, quantitative perceived agency, and non-scalar Relational Positioning - the bound relationship token, profile-sensitive contribution condition, and perspective-sensitive moral-subject gate in Section 4a are constitutive. The three route families and three identity forms govern this version; their empirical indicators, access realizations, branch criteria, and frame-selection sketches are provisional.

0.6 Judgment, not behavior

The target is the formation of judgment. A judgment may guide action, but action remains subject to weakness of will, competing motives, incapacity, coercion, habit, and environmental constraints. Nothing in the formalism entails that a detached judgment will be enacted. Conversely, behavior that conforms to a moral rule need not have been generated by moral cognition.

0.7 Typed vocabulary

The principal sorts are kept distinct throughout:

- s: the cognitive system whose active state is being modeled and from whose perspective all attributions are made;
- c: an active representational context;
- f: a frame token, understood as one active construal-organization event fixing relevant agent, object, episode, action, outcome, relevance, possibility, and local evaluative organization;
- r: a representation;
- sigma: a semantic status of a representation, such as success, relevance, or error;
- omega: a represented outcome or world-state; Delta(Omega) is the set of probability distributions over outcomes, and delta_omega is the degenerate prospect assigning probability one to omega;
- e: an outcome-evaluative token containing an outcome ordering and, where represented, outcome-level requirements;
- delta: a decision-model token fixing a frame, a frame-relative agent construal and its thin referent, a thin temporal anchor, evaluative token, represented normative perspective, requirement-source and addressee structure, set-valued object scope, action-outcome map, feasible set, unified requirement set, and priority graph;
- pi: a represented normative perspective under which a decision model treats standards as available for appraisal or guidance;
- nu: a represented source of a requirement, such as a relationship, policy, role, convention, personal commitment, or institutional rule;
- kappa: a judgment-derivation token fixing an event anchor, decision model, typed content, assessment mode, broad case, episode construal, actual support graph, and, where relevant, a violated requirement, appraisal target, local dimension, or global question;
- rho: a provenance-closed minimally sufficient derivation route realized inside the actual support graph of a frame, model, alignment, appraisal, or judgment token;
- ell: an inferential link or warrant within a derivation route;
- chi: a connected frame-model-token chain witness realized inside the judgment token's actual support graph;
- alpha: a represented execution or plan token in a context-level cross-frame action space;
- beta: a represented cross-frame action-alignment model;
- gamma: a represented cross-frame alignment model for local appraisal dimensions;
- m: a represented relationship token whose satisfaction condition may instantiate the moral-relationship schema;
- y_A: a represented agent referent that may be construed differently across frames;
- y: a represented object referent or case entity that may be construed differently across frames;
- o: a frame-relative object construal of object referent y;
- zeta: a broad represented case or occurrence token shared by outputs concerning one situation;
- epsilon: a frame-relative episode construal of broad case zeta;
- n: a node in the composed provenance graph of a connected live chain;
- n_i: an independently typed relation-bearing evaluation-input node at which the relationship, subject-binding, and object-binding paths converge;
- u0: a generic represented support unit in a composed provenance graph;
- w: a route product eligible for export, re-export, and typed import into a child route;
- x: a frame-relative agent construal of agent referent y_A whose conduct is prescribed or appraised;
- a, b: frame-relative action descriptions;

- t: a thin represented temporal anchor; frame-relative duration, granularity, and temporal description are carried by episode construals and execution tokens;
- C: a nonempty set of frame-relative action descriptions;
- u: an action-level or outcome-level requirement whose source, perspective, applicability, and possible addressee are represented;
- v: a represented priority ground capable of supporting an override;
- p: a premise or representational component;
- d: an independently identifiable active consideration;
- q: a frame-relative local appraisal dimension;
- k_rel: a represented relation-kind token used by a relation qualifier;
- Q: a thin represented overall appraisal-question token that may be shared across frames;
- z: a generic derivational target, such as a frame, decision model, alignment model, appraisal model, or judgment token;
- Γ_z , Γ_{κ} : actual functional support graphs; Γ_z is the graph of generic target z, while Γ_{κ} characterizes the anchor-identified judgment occurrence and may contain several converging realized routes, but it does not additionally individuate the judgment token;
- η_f , η_{κ} : event anchors individuating frame activations and judgment-derivation tokens;
- ϕ_{Δ} : an optional, typed, represented violation-sensitive fallback preorder over Feas_{Δ} , used only when Acc_{Δ} is empty; its model-instance rules determine severity, multiplicity, ties, and incomparability;
- z_A : a typed appraisal target, which may be an agent, action, outcome, event, object, relation, institution, or way of life;
- D1,D2: represented outcome distributions compared by the functional control-dispositional preorder $\succsim_{\text{ctrl},s,c}$;
- j: a typed judgment content.
- ψ : an active process token carrying a support unit on a live chain;
- I, J: registered intervention tokens; I indexes a bender-removal comparison and J indexes a bender-path comparison where those roles are distinguished;
- i_{ctrl} : an independently specified control-state indicator;
- g: a represented stake-architecture map from semantic statuses to represented outcome distributions;
- μ : an independently fixed measure over the pre-specified semantic contrasts in a stake-map comparison; the displayed symbol μ denotes the same sort;
- M: an actual-contribution identification model;
- ϕ : an independently specified realization-profile dimension;
- $\mathcal{R}$: an independently delimited population of representation tokens;
- $\mathcal{U}_{\text{plus}}$, $\mathcal{U}_{\text{minus}}$: respectively the positive and negative coordinates of an inherent-valence profile;
- G_{subject} : an exact ordered subject-gate configuration record;
- P_{core} : an exact token-resolved relationship-core snapshot;
- Q_{core} : the normalized payload derived from a relationship-core snapshot;
- S, S': candidate realization-profile schemes; $S^*_{\{s,K,c\}}$ is the fixed constitutively adequate scheme selected prospectively for judging system s, declared comparison class K, and context c;
- K: a declared comparison class; $K^*_{\{\kappa,c\}}$ is the class fixed for judgment token κ in c;
- h: a map between realization-profile spaces;
- E: the token-relative evaluation structure selected by RelevantEval;
- n_o : an object-binding support node;

- lambda: a relational-route witness. Some Relational Positioning witness is required for moral type. The governed route families are identity-mediated positioning, intimacy, and direct encounter; they may overlap, but at least one must be realized. The previously declared r remains a representation and is not reused as a route variable.
- iota: an independently specified identity or role-schema token that may bridge the exact represented subject and object through shared identity, contrastive identity, or complementary roles.
- V: an inherent-valence profile containing independently represented positive and negative inherent-value components.

The decision model distinguishes an agent referent y_A from its frame-relative agent construal x . First-personality is determined by whether y_A is represented as the judging system itself, not by literal identity between x and s ; third-personality is determined at the same referent level. $\text{ModeOf}(\kappa)$ is prospective or reconstructive. $\text{CaseOf}(\kappa)$ anchors the token to a broad occurrence, while $\text{EpisodeConstrualOf}(\kappa)$ records how its host frame understands that occurrence. $\text{AgentConstrual}_s(x, y_A, f, c)$ and $\text{ObjectConstrual}_s(o, y, f, c)$ distinguish referents from the agents and objects they become within frame f . A frame-relative agent, action, episode, object, or local appraisal dimension is not automatically identical to its counterpart in another frame; the relevant cross-frame bridges must be represented. Within a fixed frame and context, one active agent-construal token has at most one agent anchor; competing or uncertain reference assignments are represented by distinct construal or binding tokens rather than by one x simultaneously denoting two referents.

Reference-anchor convention. Thin agent and object referents, broad cases, temporal anchors, and appraisal questions are representational addresses within s 's active context, not thick descriptions or external metaphysical identity claims. Equality of y_A , y , z , t , or Q means that two structures are co-indexed to one anchor token. The contestable claims are the bindings from frame-relative construals to those anchors - AgentConstrual , ObjectConstrual , ConstruesEpisode , and Answers - and those binding claims must be carried by the relevant structural routes and may be undercut. If two construals are not bound to one common anchor, NDT treats them as distinct or unaligned rather than as polyphony about one entity. This is the deliberate within-context regress stop: an implementation may represent competing or uncertain binding hypotheses, but any further co-reference model would still require anchor tokens at its endpoints.

System-binding convention. Every formed model instance has exactly one modeled judging system:

$\text{JudgingSystemOf}(c)=s$ and

$\text{JudgingSystemOf}(c)=s$ and $\text{JudgingSystemOf}(c)=s' \rightarrow s=s'$.

Where an s index is suppressed, it denotes this unique system. A context that contains representations of several possible judges must either select one system at the model-instance boundary or be represented as distinct system-indexed model instances; an ambiguous binding is not an evaluable instance.

- Φ_κ : a pre-specified fine-grained token profile, potentially including content, latency, confidence, affective intensity, persistence, and appraisal profile;
- C_{contrib} : a pre-specified actual-contribution model governing redundant support and permissible intervention contingencies;
- pol : the positive or negative polarity of a generic appraisal content;
- τ : a represented cross-frame target-alignment model for actions, events, outcomes, objects, institutions, or relations;
- $\text{MoralSubjectOf}_s(\kappa, x, y_A, \chi, c)$: the route-carried represented subject under which token κ is morally typed on chain χ ;
- $\text{EmpiricallyUnresolved}(\kappa, \{H_1, \dots, H_n\}, E)$: available evidence E does not identify one provenance hypothesis among the live alternatives.

0.8 Judgment contents, identity, and actual support graphs

A bare content such as "she ought to stop" does not determine which judgment event occurred or why. NDT therefore distinguishes the content, the judgment-derivation token, the actual support graph, and the empirical evidence used to identify that graph.

$\text{SupportGraphOf}(z, c)=\Gamma_z$

$\text{SupportGraphOf}(z, c)=\Gamma$ and $\text{SupportGraphOf}(z, c)=\Gamma' \rightarrow \Gamma=\Gamma'$

Γ_z is the event-specific functional dependency structure attributed to target z in context c . It may contain several converging routes. Calling it actual does not require a literally stored graph, but it does require more than mere availability: the relevant process must have been active in formation of the token event.

0.8.1 Actual contribution under redundant support

A candidate realization-profile scheme proposes the dimensions by which materially different realizations of judgment events are represented for a declared comparison class K .

$\text{CandidateProfileScheme}(S, K, c)$

K and the scheme for K are fixed prospectively, before any individual target verdict in K is observed. Genuinely different declared comparison classes may use different schemes. A scheme may not be introduced, re-indexed, or revised after an individual verdict in order to obtain that verdict.

Methodological admissibility constrains inquiry but does not guarantee that the scheme captures the event correctly.

$\text{MethodologicallyAdmissible}(S, K, c) \leftrightarrow \text{FixedIndependently}(S, c) \text{ and } \text{ClassificationIndependent}(S, c) \text{ and } \text{DeclaredInvarianceDomain}(S, K, c)$

The target against which adequacy is assessed is functional organization, specified prior to actual contribution:

$\text{FunctionalOrganizationOf}(\kappa, K, c) := \langle \text{ActiveProcesses}(\kappa, c), \text{DependencyStructure}(\kappa, c), \text{TransitionDispositionProfile}(\kappa, K, c), \text{InvarianceConditions}(K, c) \rangle$

$\text{OrganizationRelevantDifference}(\kappa_1, \kappa_2, K, c) \leftrightarrow \text{exists } I \text{ in } \text{RegisteredInterventions}(K, c) [\text{TransitionDispositionProfile}(\kappa_1, I, K, c) \neq \text{TransitionDispositionProfile}(\kappa_2, I, K, c)]$

$\text{TracksFunctionalOrganization}(\phi, K, c)$ means that variation in dimension ϕ corresponds, at the declared resolution, to a difference in active process, dependency structure, or transition/disposition profile under the registered invariance conditions, rather than merely to an observational recoding or a change in investigator vocabulary.

$\text{OrganizationPreserving}(S, K, c) \leftrightarrow [\text{for all } \kappa_1, \kappa_2 \text{ in } K, \text{OrganizationRelevantDifference}(\kappa_1, \kappa_2, K, c) \rightarrow \text{ProfileDistinct_S}(\kappa_1, \kappa_2, K, c)] \text{ and } [\text{for all } \phi, \text{DimensionOf}(\phi, S, K, c) \rightarrow \text{TracksFunctionalOrganization}(\phi, K, c)]$

$\text{ConstitutivelyAdequate}(S, K, c) \leftrightarrow \text{OrganizationPreserving}(S, K, c)$

Methodological admissibility and constitutive adequacy are therefore distinct. A scheme can be fixed independently yet fail to preserve the target organization; conversely, adequacy is a claim about the modeled event, not a reward for pre-registration.

$\text{RealizationProfile_S}(\kappa, K, c) := \{\langle \phi, v \rangle : \text{DimensionOf}(\phi, S, K, c) \text{ and } \text{Realizes}(\kappa, \phi, v, c)\}$

$\text{FunctionallyActiveChain}(\chi, \kappa, c) \leftrightarrow \text{the frame, model, and token processes composing } \chi \text{ were functionally active in producing or sustaining } \kappa, \text{ independently of whether their support survives the selected challenge semantics}$

$\text{LiveChain}(\chi, \kappa, c) \rightarrow \text{FunctionallyActiveChain}(\chi, \kappa, c)$

$\text{ContributionVerdict_S}(\chi, \kappa, K, c) \leftrightarrow \text{FunctionallyActiveChain}(\chi, \kappa, c) \text{ and } \text{MemberOfActualMinimalSufficientSet}(\chi, \text{RealizationProfile_S}(\kappa, K, c), c) \text{ and } \text{ChainProfileDifferenceMaker_S}(\chi, \kappa, K, c)$

$\text{ContributionEquivalent}(S, S', K, c) \leftrightarrow \text{for all } \kappa \text{ in } K \text{ for all } \chi [\text{ContributionVerdict_S}(\chi, \kappa, K, c) \leftrightarrow \text{ContributionVerdict_S'}(\chi, \kappa, K, c)]$

$\text{SameDeclaredResolution}(S, S', K, c)$: S and S' represent K under one declared resolution and one comparison-relevant invariance domain

A cross-scheme bridge is specified prior to $\text{ContributionEquivalent}$. It must preserve the event organization represented by the profiles, not merely reproduce a preferred contribution verdict.

$\text{ProfileIsomorphism}(h, S, S', K, c) \leftrightarrow \text{BijectiveStructureMap}(h, \text{ProfileSpace}(S, K, c), \text{ProfileSpace}(S', K, c)) \text{ and for all } \kappa \text{ in } K [\text{h}(\text{RealizationProfile_S}(\kappa, K, c)) = \text{RealizationProfile_S'}(\kappa, K, c)]$

$\text{InterventionBasisNonempty}(K, c) \leftrightarrow \text{RegisteredInterventions}(K, c) \text{ is not empty}$

$\text{IntervenesOnChain}(I, \chi, \kappa, c)$: registered intervention I varies chain χ for token κ in c under the declared chain-intervention model

$\text{InvarianceConditionsHold}(I, K, c)$: every comparison-relevant condition declared invariant for K remains fixed under I

$\text{ChainProfileDifferenceMaker_S}(\chi, \kappa, K, c) \leftrightarrow \text{exists } I [I \text{ in } \text{RegisteredInterventions}(K, c) \text{ and } \text{IntervenesOnChain}(I, \chi, \kappa, c) \text{ and } \text{InvarianceConditionsHold}(I, K, c) \text{ and } \text{RealizationProfile_S}(\kappa, I, K, c) \neq \text{RealizationProfile_S}(\kappa, K, c)]$

$\text{InterventionEffectPreserving}(h, S, S', K, c) \leftrightarrow \text{for all } \kappa \text{ in } K \text{ for all } I \text{ in } \text{RegisteredInterventions}(K, c) [\text{InvarianceConditionsHold}(I, K, c) \rightarrow \text{h}(\text{RealizationProfile_S}(\kappa, I, K, c)) = \text{RealizationProfile_S'}(\kappa, I, K, c)]$

$\text{MinimalSufficiencyPreserving}(h, S, S', K, c) \leftrightarrow \text{for all } \kappa \text{ in } K \text{ for all } \chi [\text{MemberOfActualMinimalSufficientSet}(\chi, \text{RealizationProfile_S}(\kappa, K, c), c) \leftrightarrow \text{MemberOfActualMinimalSufficientSet}(\chi, \text{RealizationProfile_S'}(\kappa, K, c), c)]$

$\text{DifferenceMakerPreserving}(h, S, S', K, c) \leftrightarrow \text{for all } \kappa \text{ in } K \text{ for all } \chi [\text{ChainProfileDifferenceMaker_S}(\chi, \kappa, K, c) \leftrightarrow \text{ChainProfileDifferenceMaker_S'}(\chi, \kappa, K, c)]$

$\text{ContributionConvergenceBridge}(h, S, S', K, c) \leftrightarrow \text{InterventionBasisNonempty}(K, c) \text{ and } \text{SameDeclaredResolution}(S, S', K, c) \text{ and } \text{ProfileIsomorphism}(h, S, S', K, c) \text{ and } \text{InterventionEffectPreserving}(h, S, S', K, c) \text{ and } \text{MinimalSufficiencyPreserving}(h, S, S', K, c) \text{ and } \text{DifferenceMakerPreserving}(h, S, S', K, c)$

Conditional Proposition CC-1 (cross-scheme convergence): $\text{ConstitutivelyAdequate}(S, K, c) \text{ and } \text{ConstitutivelyAdequate}(S', K, c) \text{ and exists } h \text{ ContributionConvergenceBridge}(h, S, S', K, c) \rightarrow \text{ContributionEquivalent}(S, S', K, c)$

This implication is a formal guarantee under the displayed sufficient bridge conditions. $\text{FunctionallyActiveChain}$ is scheme-independent; $\text{MinimalSufficiencyPreserving}$ and $\text{DifferenceMakerPreserving}$ preserve the other two conjuncts. The nonempty intervention basis prevents the intervention clauses from being satisfied merely by an empty domain. Adequacy alone does not imply equivalence: two organization-preserving schemes may operate at different declared resolutions or partition the same organization differently, and minimal sufficiency is profile-relative.

$\text{CrossSchemeContributionInvariant}(K, c) \leftrightarrow \text{for all } S, S' [(\text{ConstitutivelyAdequate}(S, K, c) \text{ and } \text{ConstitutivelyAdequate}(S', K, c)) \rightarrow \text{ContributionEquivalent}(S, S', K, c)]$

$\text{CrossSchemeContributionInvariant}$ is a substantive model-class constraint, not a consequence of $\text{ConstitutivelyAdequate}$ by itself. An application that claims scheme-independent contribution across every adequate representation must adopt or establish this stronger condition. An application may instead fix one adequate scheme and scope its Boolean claims to that model instance. When live rival schemes remain evidentially unresolved and yield divergent contribution verdicts, the application reports $\text{EmpiricallyUnresolvedScheme}$ rather than selecting one by stipulation.

$\text{FixedConstitutiveSchemeOf_s}(K, c) = S^*_{\{s, K, c\}}$

$\text{ComparisonClassOf}(\kappa, c) = K^*_{\{\kappa, c\}}$

$\text{ConstitutiveTokenProfile}(\kappa, c) := \text{RealizationProfile_S}^*_{\{s, K^*_{\{\kappa, c\}}, c\}}(\kappa, K^*_{\{\kappa, c\}}, c)$

$\text{ActualContribution}(\chi, \kappa, c) \leftrightarrow \text{ContributionVerdict_S}^*_{\{s, K^*_{\{\kappa, c\}}, c\}}(\chi, \kappa, K^*_{\{\kappa, c\}}, c)$

$\text{RealizedIn}(\chi, \text{Gamma_}\kappa, c) \leftrightarrow \chi \text{ is a chain in } \text{Gamma_}\kappa \text{ and } \text{ActualContribution}(\chi, \kappa, c)$

Numerical token identity remains anchor-based under §0.8.2. The realization profile records how the event was realized; it does not replace JudgmentAnchorOf as the identity criterion.

$\text{EvaluableModelInstance_s}(K, c) \leftrightarrow \text{JudgingSystemOf}(c) = s \text{ and exists exactly one } S [\text{MethodologicallyAdmissible}(S, K, c) \text{ and } \text{ConstitutivelyAdequate}(S, K, c) \text{ and } \text{FixedConstitutiveSchemeOf_s}(K, c) = S]$

$\text{EvaluableFor_s}(\kappa, c) \leftrightarrow \text{EvaluableModelInstance_s}(\text{ComparisonClassOf}(\kappa, c), c)$

$\text{InvalidModelInstance}(K, c)$ when the proposed scheme is post hoc, classification-dependent, outside its declared domain, empty, or known not to preserve the target functional organization

$\text{EmpiricallyUnresolvedScheme}(K, \{S_1, \dots, S_n\}, E)$ when live candidate schemes remain evidentially unresolved and are not contribution-equivalent for the target verdicts

Outputs and every moral-type predicate are evaluated only inside an evaluable fixed model instance with the unique judging-system binding. Invalidity and unresolvedness withhold application of the Boolean predicate; neither state is a verdict that the cognitive token failed to output.

$\text{ContributionIDModel}(M) := \langle I_{\text{test}}, \Phi_{\text{obs}}, C_{\text{test}} \rangle$

$\text{IndicatorMap}(M, S, K, c)$ relates observed evidence to proposed realization dimensions

$\text{ContributionIDModel}$ and IndicatorMap remain epistemic. They may change evidential confidence, invalidity, or unresolvedness, but selecting them does not constitute the target contribution relation. A route may contribute even when removing it leaves the sentence unchanged, provided it participates in an active minimal sufficient process and changes a constitutively relevant realization feature.

An active route that is profile-silent - it changes no constitutively relevant feature of the token event - is not an actual contributor and cannot confer moral type, even when it remains available as backup. This constitutive verdict must be distinguished from empirical uncertainty about which candidate scheme preserves the event organization and which route changed it.

0.8.2 Judgment-token identity

$\text{JudgmentAnchorOf}(\kappa) = \text{eta_kappa}$

$\text{SameJudgmentToken}(\kappa_1, \kappa_2, c) \leftrightarrow \text{JudgmentAnchorOf}(\kappa_1) = \text{JudgmentAnchorOf}(\kappa_2)$

$\text{DerivationType}(\kappa) := \langle \text{Content}(\kappa), \text{BaseOf}(\kappa), \text{ModeOf}(\kappa), \text{CaseOf}(\kappa), \text{EpisodeConstrualOf}(\kappa), \text{SupportGraphOf}(\kappa, c) \rangle$

Two token occurrences may share a derivation type while remaining numerically distinct. Two tokens with identical verbal content may differ in decision model, frame, episode construal, support graph, moral status, or survival status. One token may contain several actually contributing routes inside one support graph; that is different from several token events.

Interpretation, content, source, support graph, and derivation type do not individuate a token in addition to the anchor. They characterize the anchor-identified occurrence. Changing any of them while holding the judgment-event anchor fixed describes a different attributed derivation of the same numerical token; repeating the same derivation type at a new anchor yields a numerically distinct token.

0.8.3 Frame-token identity

$\text{FrameAnchorOf}(f) = \text{eta_f}$

$\text{SameFrameToken}(f_1, f_2, c) \leftrightarrow \text{FrameAnchorOf}(f_1) = \text{FrameAnchorOf}(f_2)$

$\text{FrameTypeEquivalent}(f_1, f_2, c) \leftrightarrow \text{FrameSignature}(f_1, c) = \text{FrameSignature}(f_2, c)$

$\text{FrameDifferenceRelevant}(f_1, f_2, \kappa_1, \kappa_2, c)$: some difference in construal, relevance, possibility, evaluative accessibility, or local organization actually contributes to the compared outputs

Polyphony requires distinct frame activations and a derivationally relevant difference. Duplicating one frame under two labels does not create polyphony. Equivalent frame types may be activated more than once, but they count as polyphonic only where an actual difference contributes to the conflict.

0.8.4 Basis, routes, and chains

$\text{Basis}(z) := \text{premise-like nodes in } \Gamma_z$ $\text{Warrants}(z) := \text{link-like edges or warrant nodes in } \Gamma_z$

A route is a provenance-closed derivation subgraph of Γ_z . A connected chain joins one frame route, one decision-model route, and one judgment-token route through explicit imports. Candidate formation remains pre-defeat; output, organization, object bearing, subject matching, override, and residue are evaluated only over actually contributing live chains in Γ_κ .

0.9 Notation and semantic layers

The symbols are used consistently:

- $:=$ introduces a definition;
- $\leftrightarrow$ states a biconditional or constitutive identity;
- $\rightarrow$ states strict implication;
- $\Rightarrow$ states a defeasible or typical empirical generalization;
- $\text{Candidate}(c, \kappa)$ states that the type-specific structural bridge conditions generate derivation token κ before premise-, link-, or target-level undercutting is applied;
- $\text{LiveRoute}(\rho, z, c)$ states that route ρ for target z has available premises and available inferential links;
- $\text{DirectClear}(z, c)$ states that target z has no undefeated direct undercutter;
- $\text{Survives}(z, c)$ states that z is directly clear and has at least one live minimally sufficient route;
- $\text{LiveChain}(\chi, \kappa, c)$ states that a mutually anchored frame-model-token chain for κ survives on one connected set of routes;
- $\text{Outputs}(c, \kappa)$ states that κ is a structural candidate supported by at least one connected live chain;
- $\text{OutputsContent}(c, j)$ states that at least one output token in c has content j ;
- $\text{Exec}_\beta(C, \delta, c)$ is the set of execution tokens that alignment model β represents as realizing some action description in C under decision model δ .

No arrow is used ambiguously to mean both logical implication and cognitive production. Probability claims in the stake lift are empirical architectural claims. Minimal routes, import edges, graph reachability, and grounded defeat are defined over represented support structures; they do not imply conscious proof search.

0.10 No automatic global verdict

NDT does not assume that every context produces one context-global, all-things-considered judgment. A frame may construct an all-things-considered evaluative token internally through the closure operations defined below. If the system explicitly arbitrates among several frames, that arbitration is represented as a higher-order frame and decision model whose output has one of the ordinary judgment types. Higher-order arbitration is not iterated to mandatory closure: an arbitration token may itself conflict with lower- or equal-order tokens, and unless a further arbitration is actually represented, the resulting higher-order conflict remains polyphonic.

1. The Set-Theoretic Hierarchy

Beliefs contain normative beliefs; normative beliefs contain moral beliefs; and moral beliefs may contain a smaller subset endorsed as correct within some evaluative framework:

Beliefs $\supseteq$ NormativeBeliefs $\supseteq$ MoralBeliefs $\supseteq$ EndorsedMoralBeliefs

The nesting is structural. A normative belief has a prescriptive or appraisal form. A moral belief is a normative belief whose derivation has the additional grounding condition stated in Section 4. An endorsed moral belief is one that a further framework accepts. The innermost set does not presuppose moral realism; even an anti-realist can describe which judgment tokens a given framework endorses.

The formalism below generates tokens. Once retained, a token can be stored as a belief. At the content level, two tokens with the same sentence may occupy different subsets. A prison-avoidance token and a victim-grounded token may both have the content "do not assault her," while only the second is moral under the Moral Detachment Thesis. The set-theoretic classification therefore applies most precisely to tokened beliefs with provenance, not merely to strings of words.

In informal terms:

Normative inference. Given what is represented as the case, how relevant outcomes and standards are structured, what the represented agent can do, and whether at least one connected derivation chain survives, a prescription or appraisal can be detached.

Moral inference. A normative token is moral when a represented moral-relationship token is a provenance-bearing organizer of the token-relative evaluative structure along at least one connected live chain.

2. The Constitutive Evaluativity Thesis

"Information, it is argued, cannot supply its own energy." - Shweder (1992)

NDT does not derive evaluation from nothing. It distinguishes an ordering over representational success from an ordering over represented outcomes and from a practical ordering over feasible actions. Each transition must be separately represented.

2.1 Three ordering types

2.1.1 Semantic thin evaluative structure

A contentful representation has conditions of success, error, relevance, and correction. These conditions order representational performances. Let Σ_r be the semantic statuses available for representation r , and let $\succsim_{\text{sem}_r}$ be a preorder over those statuses.

$\text{SemTEStructure}(r,c) := \langle \Sigma_r, \succsim_{\text{sem}_r} \rangle$

$\text{SemTES}(r,c) \leftrightarrow r$ instantiates $\text{SemTEStructure}(r,c)$ in c

Accurate tracking may rank above inaccurate tracking, relevant discrimination above irrelevant noise, and successful categorization above error. This ordering does not yet rank world-outcomes or actions.

2.1.2 Outcome evaluative structure

An outcome-evaluative token ranks represented outcomes, supplies a typed preorder over uncertain prospects, and may contain outcome-level requirements. Let Ω_c be the represented outcome space, $\succsim_{out_e}$ the preorder over outcomes supplied by token e , $\succsim_{pros_e,c}$ the represented preorder over $\Delta(\Omega_c)$, and $ReqOut_e$ the set of represented outcome requirements, which may be empty.

$OutcomeEval(e,c) := \langle \Omega_c, \Delta(\Omega_c), \succsim_{out_e}, \succsim_{pros_e,c}, ReqOut_e \rangle$

For every ω_1, ω_2 in Ω_c , degenerate prospects must agree with the ordinary outcome order:

$\delta_{\omega_1} \succsim_{pros_e,c} \delta_{\omega_2} \leftrightarrow \omega_1 \succsim_{out_e} \omega_2$.

NDT does not impose expected utility, risk neutrality, stochastic dominance, or any other universal lift for nondegenerate prospects. The represented perspective must supply enough ordering structure to compare the prospects used by the formed decision model; otherwise the relevant guide is undefined.

$ReqOut_e$ does not define an exceptionless acceptable region. Section 3 imports outcome requirements into the same typed constraint system as action requirements, where either kind may be overridden by a higher live priority ground without being rendered false or inapplicable.

2.1.3 Action evaluative structure

A frame-relative decision model pulls an outcome ordering back onto a feasible action set and combines it with applicable action- and outcome-level requirements. Semantic success never directly selects an act.

semantic-status order $--g_{r,c} \rightarrow$ outcome structure $--h_{\delta} \rightarrow$ action structure

Accuracy can matter without determining whether one should warm oneself, preserve a frozen sample, recalibrate an instrument, record a reading, or remain still. The practical direction depends on the represented outcomes, constraints, and action-outcome model.

2.2 Functional control order and stochastic stake lift

Semantic correctness does not by itself explain why one represented outcome distribution matters more than another. The encoded case therefore requires an independently identifiable functional control asymmetry. The asymmetry must be represented by a pre-specified measurement model rather than an open disjunction of control tendencies.

$CtrlProfile_s(D1,c) := \langle Select_s(D1,c), Maintain_s(D1,c), Approach_s(D1,c), Restore_s(D1,c), Stabilize_s(D1,c) \rangle$

$D1 \succ=_{ctrl_s,c} D2 \leftrightarrow H_{ctrl}(CtrlProfile_s(D1,c)) \succ= H_{ctrl}(CtrlProfile_s(D2,c))$

$D1 \succ_{ctrl_s,c} D2 \leftrightarrow D1 \succ=_{ctrl_s,c} D2$ and $\neg(D2 \succ=_{ctrl_s,c} D1)$

H_{ctrl} is fixed before testing. It may be a weighted scalarization, a Pareto rule, or a declared lexicographic order, but its indicators, weights or priority order, tolerance, and invariance domain may not be selected after observing the desired classification. When indicators disagree, the chosen rule determines whether the comparison is weak, strict, or unresolved.

The control order is a theorist-level functional characterization of the system's standing dispositions. It is not objective welfare supplied by the analyst and need not be an occurrent preference token. A neutral representation feeding a separable controller is an attributed case. In an encoded case, the representation's content or functional identity is partly constituted by its role in the independently identified control profile.

Let $P^{reach_r,c}$ be the pre-specified ordered semantic contrasts reachable in the operating regime, $\mu^*_{r,c}$ an independently fixed measure over them, and $g_{r,c}$ a map from semantic statuses to represented outcome distributions. Fix $0 \leq \epsilon < 1/2$ in advance.

$g_{r,c} : \Sigma_r \rightarrow \Delta(\Omega_c)$

$StakeMap_epsilon(g,r,c|\mu^*) \leftrightarrow ArchitecturallyInstantiated(g,r,c)$ and $g : \Sigma_r \rightarrow \Delta(\Omega_c)$

and $Pr_(\sigma_1, \sigma_2) \sim \mu^* [g(\sigma_1) \succ=_{ctrl_s,c} g(\sigma_2)] \succ= 1 - \epsilon$

and $Pr_(\sigma_1, \sigma_2) \sim \mu^* [g(\sigma_1) \succ_{ctrl_s,c} g(\sigma_2)] > 0$

$StakeLift_epsilon(r,c|\mu^*) \leftrightarrow \text{exists } g \text{ StakeMap_epsilon}(g,r,c|\mu^*)$

The weak clause permits noise while requiring stable order preservation; the strict clause excludes a constant map. Without independently fixed indicators, aggregation, operating measure, tolerance, and architectural map, $StakeLift$ remains a schema rather than an identified empirical result.

2.3 Encoded, attributed, explicit, and derived evaluation

Primitive outcome-evaluative tokens enter cognition in three ways.

- Encoded outcome evaluation. $\text{EncOutcomeEval}(r,e,c)$ holds when e is realized through an architecturally instantiated stake map whose consumer relation is constitutive of r 's content or functional identity.
- Attributed outcome evaluation. $\text{AttrOutcomeEval}(r,e,c)$ holds when a separable downstream evaluator applies e to an otherwise neutral representation.
- Explicit evaluation. $\text{ExplicitEval}(e,c)$ holds when a goal, desire, command, policy, ordering, or standard is active as a standalone state.

$$\text{EncOutcomeEval}(r,e,c) \leftrightarrow \text{OutcomeEval}(e,c) \wedge \exists g[\text{StakeMap}_\varepsilon(g,r,c|\mu^*_r,c)]$$

$$\wedge \text{Realizes}(e,g,r,c) \wedge \text{ConstitutiveConsumerRole}(g,e,r,c)]$$

$\text{Realizes}(e,g,r,c)$ means that e is the outcome-evaluative token whose outcome and prospect ordering realizes the stake map g for r in c . $\text{ConstitutiveConsumerRole}(g,e,r,c)$ means that r 's standing role in the g -to- e consumer organization partly individuates r 's represented content or functional identity; a separable downstream evaluator does not satisfy this predicate. These are typed realization and consumer-role predicates, not evidence that any real representation instantiates them.

This is a classification condition, not the empirical claim that encoded cases are common. The primitive base is:

$$\text{PrimitiveEvalBase}(c) := \{e : \exists r \text{ EncOutcomeEval}(r,e,c) \vee \exists r \text{ AttrOutcomeEval}(r,e,c) \vee \text{ExplicitEval}(e,c)\}$$

Human reasoning can combine several primitive tokens. Let A be the pre-specified family of aggregation, constraint-combination, and priority operations available to the system. For frame f :

$$\text{EvalClosure}(f,c) := \text{Cl}_A(\{e \in \text{PrimitiveEvalBase}(c) : \text{AccessibleTo}(e,f,c)\})$$

Every derived token must retain a provenance map $\text{Prov}(e)$ to the primitive and previously derived tokens from which it was constructed. Aggregate, lexicographic, and constraint-sensitive structures therefore enter as traceable derived tokens rather than untyped additions.

2.4 Strong, weak, and practical versions

Acknowledged strong semantic-CET argument, not an NDT assumption. One external philosophical argument holds that contentful representation constitutively involves correctness, relevance, or error conditions and therefore instantiates SemTES. NDT does not adopt that universal premise: no definition, candidate rule, output rule, moral-type rule, or internal lemma below depends on every representation instantiating SemTES.

Weak evolutionary or learning rationale, also non-governing. Selection and learning may help explain an association between semantic differentiation and control-sensitive organization. That rationale does not establish pervasiveness, adaptation, encodedness, or any particular causal history. The registered H14 states only the association it tests.

Practical CET must be stated without defining its antecedent by its conclusion. $\text{ControlRecruited}(r,c)$ requires evidence that r is consumed by a standing policy-selection architecture and that an intervention on r changes an independently identified policy state, action model, or control disposition. The control-state indicator, measurement rule, and operating range must be fixed before the target output is observed and may not be defined by that same output or behavior. A neutral sensor coupled to a separable controller may satisfy this condition, so control recruitment does not entail encodedness.

$$\text{IndependentControlIndicator}(i_ctrl,r,c) \leftrightarrow \text{PreSpecified}(i_ctrl,c) \wedge \text{TracksStandingControlState}(i_ctrl,r,c) \wedge \neg \text{DefinedByTargetOutput}(i_ctrl,c)$$

$$\text{PolicySensitiveToIntervention}(r,c) \leftrightarrow \exists i_ctrl[\text{IndependentControlIndicator}(i_ctrl,r,c) \wedge \text{ChangesUnderIntervention}(i_ctrl,r,c)]$$

$$\text{ControlRecruited}(r,c) \leftrightarrow \text{ConsumedByStandingControl}(r,c) \wedge \text{PolicySensitiveToIntervention}(r,c)$$

$\text{CorrectOrRelevant}(r,o,c)$ means that r is represented in c as correct, successful, or relevant with respect to represented object or state o . This declared predicate replaces the previously unreduced abbreviation CR. $\text{Active}(r,c)$ means that r participates functionally in the modeled episode rather than merely being available.

The first practical claim concerns action relevance. It is the open empirical generalization H1b, not a constitutive axiom:

$$\text{CorrectOrRelevant}(r,o,c) \wedge \text{Active}(r,c) \wedge \text{ControlRecruited}(r,c) \\ \Rightarrow \exists g,e[\text{StakeMap}_\varepsilon(g,r,c|\mu^*) \wedge \text{OutcomeEval}(e,c) \wedge \text{Realizes}(e,g,r,c)]$$

Control recruitment alone entails neither encodedness, stable order preservation, nor action guidance; each conclusion requires its own antecedents and evidence.

The distinctive encoded claims are separately registered over one independently delimited population $\mathcal{R}$ of representation tokens in context c , with sampling measure $v_{\mathcal{R}}$ fixed in advance:

$$\text{EncodedRate}(\mathcal{R},c) := \Pr_{(r \sim v_{\mathcal{R}})}[\exists e \text{ EncOutcomeEval}(r,e,c)]$$

$$\text{EncodedExistence}(\mathcal{R},c) \leftrightarrow \exists r \in \mathcal{R} \exists e \text{ EncOutcomeEval}(r,e,c)$$

$$\text{EncodedNonNegligible}_{\mathcal{G}}(\mathcal{R},c) \leftrightarrow \text{NonNegligibleBy}(\mathcal{G},\text{EncodedRate}(\mathcal{R},c),\mathcal{R},c)$$

$\mathcal{G}$ is future empirical governance fixing the sampling frame, measurement rule, uncertainty treatment, and meaning of non-negligible before data are inspected. NDT fixes no universal numerical threshold. H15a concerns existence; H15b concerns non-negligible occurrence at the same representation-token unit. Neither follows from H1, H1b, or a person-level existential case. Attributed cases can satisfy control recruitment and H1b while failing EncOutcomeEval.

2.5 Formal expression and barriers

$$\text{SemTES}(r,c) \rightarrow \neg \text{SemEvalFree}(r,c)$$

The previously unreduced DescriptiveCIP predicate is eliminated. An explicit or attributed evaluator enters PrimitiveEvalBase directly under §2.3; no second name is needed for that route.

SemTES alone does not entail OutcomeEval, and OutcomeEval alone does not entail an action or prospect ordering adequate for a formed decision problem. Control recruitment is independently operationalized; H1b remains defeasible; and encoded existence and non-negligibility are separate empirical commitments. An evolutionary story is optional rationale only: not every feature of an evolved architecture was selected for its current role, and learned byproducts or spandrels may contribute to individual tokens.

3. The Normative Detachment Thesis

NDT states the computational form by which traceable evaluative structures, frame-relative decision models, unified constraints, priority relations, appraisal structures, and surviving derivation routes generate normative judgment tokens.

3.1 Decision-model tokens, normative perspectives, and object scope

A decision-model token delta fixes one local practical problem. It records a frame f , frame-relative agent construal x , thin agent referent y_A , thin temporal anchor t , outcome-evaluative token e , represented normative perspective π , nonempty feasible set, action-outcome map, action- and outcome-level requirements with represented sources and possible addressees, an operative priority graph, a set-valued object scope, and a provenance basis.

$$\text{FrameOf}(\delta) = f \text{ AgentOf}(\delta) = x \text{ AgentReferentOf}(\delta) = y_A$$

$$\text{TimeAnchorOf}(\delta) = t \text{ OrdOf}(\delta) = e \text{ Basis}(\delta) = \Pi\delta$$

$$\text{PerspectiveOf}(\delta) = \pi \text{ SourceOf}(u,\delta) = v$$

PerspectiveRepresented_s(π,δ,c): π is the standpoint under which δ treats standards as available for appraisal or guidance

SourceRepresented_s(v,u,δ,c): v is represented as the source of requirement u in δ

AddressedTo_s(u,y_A,x,t,f,c): u is represented as directed to agent referent y_A under construal x at t in f

$$\text{NormativeScope}_s(\delta,c) := \langle \text{PerspectiveOf}(\delta), \text{SourceOf}(\cdot,\delta),$$

$$\text{AddressedTo}_s(\cdot, \text{AgentReferentOf}(\delta), \text{AgentOf}(\delta), \text{TimeAnchorOf}(\delta), \text{FrameOf}(\delta), c) \rangle$$

$$\text{AgentConstrual}_s(x,y_A,f,c)$$

$$\text{AgentConstrual}_s(x,y_A,f,c) \wedge \text{AgentConstrual}_s(x,y'_A,f,c) \rightarrow y_A = y'_A$$

$$\text{AgentAnchorOf}_s(x,f,c) = y_A \leftrightarrow \text{AgentConstrual}_s(x,y_A,f,c)$$

$$\text{AgentAnchorBinding}_s(\delta,y_A,c) \leftrightarrow \text{AgentReferentOf}(\delta) = y_A \wedge \text{AgentConstrual}_s(\text{AgentOf}(\delta),y_A,\text{FrameOf}(\delta),c)$$

AgentBindingClaim_s(δ,y_A,c): represented premise token whose satisfaction condition is AgentAnchorBinding_s(δ,y_A,c)

$$\text{CarriesAgentBinding}_s(\rho,\delta,y_A,c) \leftrightarrow \text{AgentBindingClaim}_s(\delta,y_A,c) \in \text{Premises}(\rho) \wedge \text{AgentAnchorBinding}_s(\delta,y_A,c)$$

$$\text{Req}^{\text{out}}_{\delta} := \text{Import}_{\delta}(\text{ReqOut_OrdOf}(\delta))$$

$$\text{Req}^{\text{act}}_{\delta} := \text{Req}^{\text{act}}_{\delta} \cup \text{Req}^{\text{out}}_{\delta}$$

$$\text{PriorityDomain}(\delta) := \text{Req}^{\text{act}}_{\delta} \cup \text{Grounds}_{\delta}$$

$\text{DecisionEval}(\delta) := \langle \text{OrdOf}(\delta), \text{Req}^*_{\delta}, \text{NormativeScope}_s(\delta, c), \triangleright_{\delta} \rangle$

$\text{DecisionClosure}(f, c) := \text{Cl}_B(\{\text{accessible outcome tokens, requirements, normative perspectives, source and addressee materials, and priority materials in } f\})$, with provenance defined

Let h_{δ} map feasible actions to represented outcome distributions. The action preorder is the pullback of the represented prospect preorder, not an illicit application of the outcome preorder to distributions.

$h_{\delta} : \text{Feas}_{\delta} \rightarrow \Delta(\Omega_c)$

$a \succcurlyeq_{\delta} b \leftrightarrow h_{\delta}(a) \succcurlyeq_{\text{pros-OrdOf}(\delta), c} h_{\delta}(b)$

$\text{Best}_{\delta} := \{a \in \text{Feas}_{\delta} : \neg \exists b \in \text{Feas}_{\delta} \text{ such that } b \succ_{\delta} a\}$

$\text{OutcomeAbout}_s(h_{\delta}(a), o, y, \delta, c)$: the represented outcome distribution for a concerns object construal o of referent y

$\text{RequirementAbout}_s(u, o, y, \delta, c)$: requirement u is represented as concerning object construal o of referent y

$\text{ActionTreats}_s(a, o, y, \delta, c)$: action description a is represented as treating object construal o of referent y

$\text{ObjectScope}_s(\delta, c) := \{\langle o, y \rangle : \text{ObjectConstrual}_s(o, y, \text{FrameOf}(\delta), c) \wedge [\exists a \in \text{Feas}_{\delta} (\text{OutcomeAbout}_s(h_{\delta}(a), o, y, \delta, c) \vee \text{ActionTreats}_s(a, o, y, \delta, c)) \vee \exists u \in \text{Req}^*_{\delta} \text{RequirementAbout}_s(u, o, y, \delta, c)]\}$

$\text{ObjectAnchorBinding}_s(\delta, o, y, c) \leftrightarrow \langle o, y \rangle \in \text{ObjectScope}_s(\delta, c) \wedge \text{ObjectConstrual}_s(o, y, \text{FrameOf}(\delta), c)$

$\text{ObjectBindingClaim}_s(\delta, o, y, c)$: represented premise token whose satisfaction condition is $\text{ObjectAnchorBinding}_s(\delta, o, y, c)$

$\text{CarriesObjectBinding}_s(\rho, \delta, o, y, c) \leftrightarrow \text{ObjectBindingClaim}_s(\delta, o, y, c) \in \text{Premises}(\rho) \wedge \text{ObjectAnchorBinding}_s(\delta, o, y, c)$

$\text{CarriesObjectScope}_s(\rho, \delta, c) \leftrightarrow \forall \langle o, y \rangle \in \text{ObjectScope}_s(\delta, c) \text{ CarriesObjectBinding}_s(\rho, \delta, o, y, c)$

At token level Feas_{δ} is finite. A continuous implementation requires the ordinary compactness and continuity conditions needed for an undominated set to exist.

$\text{FallbackOf}(\delta) \in \text{Optional}(\text{Preorder}(\text{Feas}_{\delta}))$

$\text{FallbackOf}(\delta) = \varphi_{\delta}$ when a fallback preorder is represented; $\text{SupportGraphOf}(\delta, c) = \Gamma_{\delta}$

$\text{ModelRouteRealized}(\rho, \delta, c) \leftrightarrow \text{Route}(\rho, \delta, c) \wedge \rho \subseteq \Gamma_{\delta}$

The fallback structure may be unused in ordinary cases, but a formed model must contain or explicitly omit it. Its represented rules determine severity, multiplicity, ties, and incomparability. If Acc_{δ} is empty and no fallback preorder is represented, or if that preorder does not define a maximal set for the feasible case, the model does not generate comparative guidance merely by reverting to the outcome ordering.

3.2 Applicable constraints, priority, override, acceptability, and guidance

Action- and outcome-level constraints use one typed system. Perspective is distinct from source and addressee: an observer may apply a communal or role-relative standard to an agent who does not endorse it, but blame additionally requires that the agent be represented as answerable to the relevant requirement. Philosophical labels do not settle any of these facts.

$\text{PerspectiveLicenses}_s(\pi, u, \delta, c)$: under π , requirement u is available for appraisal or guidance in δ

$\text{ActiveRequirement}_s(u, \delta, c) \leftrightarrow u \in \text{Req}^*_{\delta} \wedge \text{PerspectiveLicenses}_s(\text{PerspectiveOf}(\delta), u, \delta, c)$

$\text{ApplicableRequirement}_s(u, a, \delta, c) \leftrightarrow \text{ActiveRequirement}_s(u, \delta, c) \wedge \text{Applies}(u, a, \delta, c)$

$\text{Satisfies}_{\delta}(a, u, c) \leftrightarrow [\text{Kind}(u) = \text{action} \wedge \text{SatisfiesAct}(a, u, \delta, c)] \vee [\text{Kind}(u) = \text{outcome} \wedge \text{SatisfiesOut}(h_{\delta}(a), u, \delta, c)]$

$\text{SatisfactionDeterminate}_{\delta}(a, u, c)$ means that the formed evaluable model represents enough information to decide $\text{Satisfies}_{\delta}(a, u, c)$. Missing, unresolved, or merely unmeasured satisfaction evidence is not nonsatisfaction.

$\text{Violates}(a, u, \delta, c) \leftrightarrow \text{ApplicableRequirement}_s(u, a, \delta, c) \wedge \text{SatisfactionDeterminate}_{\delta}(a, u, c) \wedge \neg \text{Satisfies}_{\delta}(a, u, c)$

The raw priority graph may be incomplete, but a formed model requires a finite, irreflexive, acyclic operative graph. Its transitive closure supplies comparisons.

$\text{PriorityWellFormed}(\delta) \leftrightarrow \text{Finite}(\text{PriorityDomain}(\delta)) \wedge \text{Irreflexive}(\triangleright_{\delta}) \wedge \text{Acyclic}(\triangleright_{\delta})$

$\triangleright^+_{\delta} := \text{transitive closure of } \triangleright_{\delta}$

$\text{LicensedOverride}(u, a, \delta, c) \leftrightarrow \text{Violates}(a, u, \delta, c) \wedge \exists v [\text{OverrideGround}(v, u, a, \delta, c) \wedge v \triangleright^+_{\delta} u]$

LicensedOverride is structural: it says that the formed model contains a qualifying violation, ground, and priority path.

$\text{OperativeOverrideOn}(u, a, v, w, \chi, \kappa, c)$ is separate and requires LicensedOverride , a ground node v and warrant w on the same RealizedLiveChain χ , and an actually contributing path from them to $\text{RelevantEval}(\kappa)$. Structural licensing can therefore be present while no override is operative in an output.

$\text{Disqualifies}(u,a,\delta,c) \leftrightarrow \text{Violates}(a,u,\delta,c) \wedge \neg \text{LicensedOverride}(u,a,\delta,c)$

$\text{Acc}_\delta := \{a \in \text{Feas}_\delta : \neg \exists u \in \text{Req}^*_\delta \text{Disqualifies}(u,a,\delta,c)\}$

When acceptable acts exist, guidance maximizes the outcome pullback inside that region. When every feasible act is unacceptable, the decision model must use a represented violation-sensitive fallback preorder φ_δ . NDT does not constitutively require one universal arbitration rule, but it forbids the constraints from disappearing at the moment they become unavoidable.

$\text{ViolationProfile}_\delta(a)$:= the vector of non-overridden violations by descending represented priority, with any represented severity or multiplicity information

$a \succcurlyeq_{\text{fallback}_\delta} b \leftrightarrow \text{FallbackOf}(\delta) = \varphi_\delta$ and φ_δ ranks a no worse than b using $\text{ViolationProfile}_\delta$ and, where tied, $\succcurlyeq_\delta$

$\text{Guide}_\delta := \text{Max}_{\succcurlyeq_\delta}(\text{Acc}_\delta)$, if $\text{Acc}_\delta \neq \emptyset$

$\text{Guide}_\delta := \text{Max}_{\succcurlyeq_{\text{fallback}_\delta}}(\text{Feas}_\delta)$, if $\text{Acc}_\delta = \emptyset$, $\text{FallbackOf}(\delta) = \varphi_\delta$, and that represented preorder defines the maximal set

Guide_δ undefined, if $\text{Acc}_\delta = \emptyset$ and no applicable represented fallback structure is defined

$\text{Nontrivial}_\delta \leftrightarrow \emptyset \neq \text{Guide}_\delta \subseteq \text{Feas}_\delta$

$\text{WrongRelative}(a,\delta) \leftrightarrow a \in \text{Feas}_\delta \wedge a \notin \text{Acc}_\delta$

WrongRelative and every Wrong output are relative to the represented model δ . They do not assert objective wrongness or justification.

$\text{UnavoidableViolation}(\delta) \leftrightarrow \text{Acc}_\delta = \emptyset$

$\text{TragicChoice}(a,\delta) \leftrightarrow a \in \text{Guide}_\delta \wedge \text{WrongRelative}(a,\delta)$

$\text{StructuralResidue}(a,u,\delta,c) \leftrightarrow \text{Violates}(a,u,\delta,c) \wedge [\text{LicensedOverride}(u,a,\delta,c) \vee \text{TragicChoice}(a,\delta)]$

$\text{OutputResidue}(a,u,\kappa,c) \leftrightarrow \text{Outputs}(c,\kappa) \wedge \text{BaseOf}(\kappa) = \delta \wedge \text{Violates}(a,u,\delta,c) \wedge [\exists v,w,\chi \text{OperativeOverrideOn}(u,a,v,w,\chi,\kappa,c) \vee \text{TragicChoice}(a,\delta)]$

A lexicographic fallback that minimizes the highest-priority violations before considering lower-priority violations and outcomes is one admissible restricted family, not a constitutive axiom. The exact fallback is empirically testable and must be represented with traceable provenance.

Nontrivial_δ governs the formal OughtChoose output only. It deliberately withholds that output for singleton feasible sets and all-tied guides; it does not settle every ordinary-language use of ought, requirement, preference, permission, authorization, exception, or protected choice. Those additional grammar types remain outside the current six-output grammar.

3.3 Judgment grammar and appraisal structures

Normative judgment contents are typed as follows:

$j ::= \text{OughtChoose}(x,C,t) \mid \text{Violation}(x,a,u,t) \mid \text{Wrong}(x,a,t)$

$\mid \text{Blameworthy}(x,a,t) \mid \text{Praiseworthy}(x,a,t) \mid \text{Appraise}(z_A,q,\text{pol},t)$

- $\text{OughtChoose}(x,C,t)$ is prospective comparative guidance.
- $\text{Violation}(x,a,u,t)$ is requirement-relative and may coexist with guidance, override, wrongness, or praise.
- $\text{Wrong}(x,a,t)$ places an act outside the overall acceptable region of the indexed model.
- $\text{Blameworthy}(x,a,t)$ adds answerability and requirement-specific responsibility to overall act wrongness.
- $\text{Praiseworthy}(x,a,t)$ is a specialized favorable agent-act appraisal that attributes relevant responsibility or credit. Admiration and favorable appraisal without credit use Appraise rather than this type.
- $\text{Appraise}(z_A,q,\text{pol},t)$ is a generic positive or negative appraisal of a formed target z_A . It belongs to normative cognition whether or not it receives moral type.

The generic form represents negative aspectual criticism and evaluation of outcomes, events, objects, relations, institutions, or ways of life. It does not create agentless morality. A generic appraisal receives moral type only when its actual derivation contains a represented moral subject, explicit or implicit, whose relationship organizes the operative appraisal.

$\text{TargetSort}(z_A)$ in $\{\text{agent}, \text{action}, \text{agent-act}, \text{outcome}, \text{event}, \text{object}, \text{relation}, \text{institution}, \text{way-of-life}\}$

$\text{AgentActTarget}(x,a)$: a typed target constructor of sort agent-act

pol in {positive,negative}

AppStruct_delta(q):=<X_delta,q, >=app_delta,q, PosRegion_delta,q, NegRegion_delta,q, LocalScope_delta(q), LocalTarget_delta(q), Prov_delta,q>

Profile_{delta,q}(z_A,t,c) in X_delta,q

AppraisesPositively(z_A,delta,q,t,c) ↔ Profile_{delta,q}(z_A,t,c) in PosRegion_delta,q

AppraisesNegatively(z_A,delta,q,t,c) ↔ Profile_{delta,q}(z_A,t,c) in NegRegion_delta,q

Exemplary(x,a,t,delta,q,c) ↔ AppraisesPositively(AgentActTarget(x,a),delta,q,t,c)

LocalScope_delta(q) in {aspectual,overall}; LocalTarget_delta(q) in TargetSorts

AppClosure(f,c):=Cl_C({appraisal-relevant materials accessible to f}), with provenance defined

Thin overall question tokens Q remain distinct from local dimensions q. Q identifies a represented cross-frame question, not a universal value. Wrong and Blameworthy answer overall questions by type; Appraise and Praiseworthy answer Q only where a represented Answers_delta bridge exists.

QuestionStruct(Q):=<QuestionTarget(Q),QuestionKey(Q)>

Answers_delta(q,Q,c): local dimension q is represented as answering Q

3.4 Structural formation and assessment modes

Candidate generation is pre-defeat. Structural formation requires coherent and traceable resources, represented agent and object bindings, normative perspective and source structure, and formed appraisal targets. It does not require that any challenge has already been defeated.

The formation rules begin after upstream representations, anchors, resources, and support structures are supplied to the model instance. They conditionally classify and connect a candidate derivation; they are not an end-to-end temporal generator from raw sensation, memory, or unstructured cognition and do not claim prospective online dynamics.

CurrentFrameFormed(f,c) ↔ InternallyCoherent(f,c) and Represented(f,c) and DecisionResources(f,c) and exists rho Route(rho,f,c)

ReconstructionFrameFormed_s(f,x,t,c) ↔ ReconstructionCoherent_s(f,x,t,c) and ReconstructionRepresented_s(f,x,t,c) and DecisionResources(f,c) and exists rho Route(rho,f,c)

NormativeScopeFormed_s(delta,c) ↔ PerspectiveFormed_s(delta,c) and AddressStructureRepresented_s(delta,c)

ModelFormed(c,delta) ↔ DecisionEval(delta) in DecisionClosure(FrameOf(delta),c) and AgentAnchorBinding_s(delta,AgentReferentOf(delta),c) and NormativeScopeFormed_s(delta,c)

and MapRepresented(h_delta,FrameOf(delta),c) and FeasRepresented(Feas_delta,FrameOf(delta),c) and ConstraintsRepresented(Req*_delta,FrameOf(delta),c) and PriorityWellFormed(delta)

and exists rho [Route(rho,delta,c) and CarriesAgentBinding_s(rho,delta,AgentReferentOf(delta),c) and CarriesNormativeScope_s(rho,delta,c) and CarriesObjectScope_s(rho,delta,c)]

ProspectiveFormed(c,delta) ↔ CurrentFrameFormed(FrameOf(delta),c) and CurrentDecisionModel(delta,c) and ModelFormed(c,delta)

ReconstructionFormed_s(c,delta,x,t) ↔ Reconstructs_s(delta,x,t,c) and AgentOf(delta)=x and TimeAnchorOf(delta)=t and ReconstructionFrameFormed_s(FrameOf(delta),x,t,c) and ModelFormed(c,delta)

ModeOf(kappa) in {prospective,reconstructive}

AssessmentModeFormed_s(c,kappa,delta,x,t) ↔ [ModeOf(kappa)=prospective and ProspectiveFormed(c,delta) and AgentOf(delta)=x and TimeAnchorOf(delta)=t] or [ModeOf(kappa)=reconstructive and ReconstructionFormed_s(c,delta,x,t)]

AssessmentTargetFormed_s(c,kappa,delta,z_A,t) ↔ AppraisalTargetOf(kappa)=z_A and TargetSort(z_A)=LocalTarget_delta(LocalDimensionOf(kappa)) and TargetBoundInFrame_s(z_A,FrameOf(delta),CaseOf(kappa),c) and TimeAnchorOf(delta)=t

CaseEpisodeFormed_s(kappa,c) ↔

ConstruesEpisode_s(FrameOf(BaseOf(kappa)),EpisodeConstrualOf(kappa),CaseOf(kappa),c)

DerivationFormed(c,kappa) ↔ CaseEpisodeFormed_s(kappa,c) and exists rho Route(rho,kappa,c)

AssessmentTargetFormed is target-sorted. Agent-act profiles use AgentActTarget(x,a); event, outcome, object, relation, institution, and way-of-life targets require their own route-carried bindings. Surface form does not determine whether an agent is represented in the derivation.

3.5 Typed candidate generation

ProspectiveCandidate(c,kappa) $\leftrightarrow$ Content(kappa)=OughtChoose(x,C,t) and BaseOf(kappa)=delta and ModeOf(kappa)=prospective and AgentOf(delta)=x and TimeAnchorOf(delta)=t and ProspectiveFormed(c,delta) and DerivationFormed(c,kappa) and C=Guide_delta and Nontrivial_delta

Scenario(x,a,t,c) records actual, prospective, or counterfactual performance under construal x. Answerability and responsibility are requirement-specific; they are not inferred merely from an overall wrongness label.

AnswerableTo_s(y_A,x,u,t,f,c) $\leftrightarrow$ AddressedTo_s(u,y_A,x,t,f,c) and StandingAppliesTo_s(u,y_A,x,t,f,c)

ResponsibleCapacityFor_s(y_A,x,a,g,t,delta,c): s attributes control, relevant epistemic access, and opportunity to y_A-under-x with respect to the requirement or appraisal ground g

CreditsFor_s(y_A,x,a,q,t,delta,c): the favorable agent-act appraisal on dimension q attributes relevant responsibility or credit for a to y_A-under-x, rather than merely admiring or favorably appraising it

ViolationCandidate(c,kappa_v) $\leftrightarrow$ Content(kappa_v)=Violation(x,a,u,t) and BaseOf(kappa_v)=delta and AssessmentModeFormed_s(c,kappa_v,delta,x,t) and DerivationFormed(c,kappa_v) and Scenario(x,a,t,c) and Violates(a,u,delta,c)

WrongCandidate(c,kappa_w) $\leftrightarrow$ Content(kappa_w)=Wrong(x,a,t) and BaseOf(kappa_w)=delta and AssessmentModeFormed_s(c,kappa_w,delta,x,t) and DerivationFormed(c,kappa_w) and Scenario(x,a,t,c) and WrongRelative(a,delta) and QuestionOf(kappa_w)=Q and OverallActQuestion(Q,delta,c)

BlameCandidate(c,kappa_b) $\leftrightarrow$ Content(kappa_b)=Blameworthy(x,a,t) and BaseOf(kappa_b)=delta and AssessmentModeFormed_s(c,kappa_b,delta,x,t) and DerivationFormed(c,kappa_b) and Scenario(x,a,t,c) and WrongRelative(a,delta)

and exists u[Disqualifies(u,a,delta,c) and AnswerableTo_s(AgentReferentOf(delta),x,u,t,FrameOf(delta),c) and ResponsibleCapacityFor_s(AgentReferentOf(delta),x,a,u,t,delta,c)] and QuestionOf(kappa_b)=Q and OverallAgentQuestion(Q,delta,c)

PraiseCandidate(c,kappa_p) $\leftrightarrow$ Content(kappa_p)=Praiseworthy(x,a,t) and BaseOf(kappa_p)=delta and AssessmentModeFormed_s(c,kappa_p,delta,x,t) and DerivationFormed(c,kappa_p) and Scenario(x,a,t,c)

and LocalDimensionOf(kappa_p)=q and AppStruct_delta(q) in AppClosure(FrameOf(delta),c) and Exemplary(x,a,t,delta,q,c) and ResponsibleCapacityFor_s(AgentReferentOf(delta),x,a,q,t,delta,c) and CreditsFor_s(AgentReferentOf(delta),x,a,q,t,delta,c)

GenericAppraisalCandidate(c,kappa_a) $\leftrightarrow$ Content(kappa_a)=Appraise(z_A,q,pol,t) and BaseOf(kappa_a)=delta and AssessmentTargetFormed_s(c,kappa_a,delta,z_A,t) and DerivationFormed(c,kappa_a)

and AppStruct_delta(q) in AppClosure(FrameOf(delta),c) and [(pol=positive and AppraisesPositively(z_A,delta,q,t,c)) or (pol=negative and AppraisesNegatively(z_A,delta,q,t,c))]

Candidate(c,kappa) $\leftrightarrow$ ProspectiveCandidate or ViolationCandidate or WrongCandidate or BlameCandidate or PraiseCandidate or GenericAppraisalCandidate

GenericAppraisalCandidate is a normative type condition only. It does not imply that the appraisal is moral, that an agent is represented, or that FirstPersonal or ThirdPersonal applies.

3.6 Actual support graphs, connected survival, and output

Candidate is genuinely pre-defeat. Final output is evaluated against the token's event-specific support graph rather than every route that could be reconstructed from a broad basis.

ActualChainSet(kappa,c):={chi: RealizedIn(chi,SupportGraphOf(kappa,c),c)}

RealizedLiveChain(chi,kappa,c) $\leftrightarrow$ LiveChain(chi,kappa,c) and ActualContribution(chi,kappa,c)

Outputs(c,kappa) $\leftrightarrow$ EvaluableFor_s(kappa,c) and Candidate(c,kappa) and exists chi RealizedLiveChain(chi,kappa,c)

OperativeEvaluativeSourceOn_s(e,chi,kappa,c) $\leftrightarrow$ e belongs to the primitive or provenance-preserving derived evaluative base of BaseOf(kappa), e is carried on chi, and e contributes to RelevantEval(kappa) on chi.

Bounded evaluator-source lemma ES-1. For every kappa admitted by the current six-output grammar, Outputs(c,kappa) $\rightarrow$ exists e,chi[OperativeEvaluativeSourceOn_s(e,chi,kappa,c)]. This is an internal, grammar-relative consequence of

ModelFormed, the type-specific candidate rules, RealizedLiveChain, and actual contribution. It is not an external exhaustiveness theorem, does not make strong or universal CET an NDT assumption, and does not require encoded CET in every token: *e* may be explicit, attributed, encoded, or provenance-preservingly derived.

Several actually contributing live chains may converge in one token event. A moral and a nonmoral route can both be actual when each participates in an active minimal sufficient process for a constitutively relevant token feature. A merely possible, inactive, or profile-silent route does not confer output, moral type, override, or residue.

A chain can be functionally active yet fail LiveChain because an admitted challenge defeats or leaves undecided a premise, link, or target. Such a chain records causal or functional occurrence without route survival; it does not satisfy RealizedLiveChain and cannot by itself support Outputs.

Outputs is not constituted by an investigator's choice of intervention family, observable feature set, or causal-identification model. A formal application nevertheless evaluates the token inside a fixed constitutive model instance, including a declared realization-profile scheme. Candidate schemes are rival hypotheses about the event's functional organization; selecting one does not make its verdict true. If methodologically admissible candidate schemes remain verdict-divergent and available evidence does not identify an organization-preserving scheme or contribution-equivalent class, the empirical classification is unresolved rather than negative.

3.7 Explicit, encoded, mixed, and override cases

The explicit case uses an evaluative token supplied by a tokened goal, command, or policy. The encoded case uses a token supplied by a representation whose stake-sensitive consumer role is constitutive. Both feed the same decision-model architecture. Mixed cases are constructed through closure operations with traceable provenance.

A priority ground that licenses override is part of the structural candidate only when it is represented in delta. It is operative in an output only if a model route realized in the token's actual support graph and contributing to the output contains both the ground and the warrant connecting it to the override. Evidence that the ground is false, unreliable, or inapplicable undercuts that premise or link; it does not retroactively change the meaning of LicensedOverride.

3.8 Tie, multiplicity, tragedy, plurality, and conflict

Choice sets prevent multiple co-best acts from becoming several individually binding oughts. A nontrivial guide is a nonempty proper subset of the feasible set; its members are jointly offered as the guided choice set.

Tragic choice occurs when Acc_δ is empty but a represented fallback preorder selects one or more feasible acts. Such acts remain wrong relative to the model and may retain residue. Because the fallback remains sensitive to the priority and violation structure, NDT can represent "violate the lesser requirement" rather than silently switching to outcome maximization.

Intra-model tragedy is distinct from cross-frame polyphony. One decision model may guide an act while also representing it as wrong; polyphony requires distinct frame activations, relevant frame differences, alignment, and conflicting outputs.

3.9 Partitioning

The deductive part of NDT begins once a frame, agent referent and agent construal, object referent and object construal, broad case and episode construal, thin temporal anchor, represented normative perspective, requirement-source and addressee structure, set-valued object scope, decision model, unified requirements, priority graph, action-outcome map, feasible set, local appraisal structure, global question bridge where applicable, structural routes, action or dimension alignment model where needed, and grounded attack graph are fixed. The empirical part concerns how these structures are learned, represented, and revised, and how closely actual cognition approximates the formal relations.

4. The Moral Detachment Thesis

The MDT describes moral cognition as a structurally defined subset of normative cognition.

MDT is a distinct constitutive proposal applied after normative output. It is not derived from semantic CET, encoded CET, or the general evaluative-prerequisite lemma. Its external adequacy as an account of the moral kind remains open under the registered hypotheses and identification programs.

MDT. An output token is moral exactly when one or more represented moral-relationship tokens organize the token-relative evaluative structure along at least one connected live chain to that indexed derivation.

The word organize is load-bearing. Mere co-occurrence, thematic relevance, or reference to the same people and actions is insufficient. A punishment-avoidance token does not become moral merely because a victim is also represented as a moral object. A relationship representation must occur in a connected live chain and have a directed provenance path to the evaluative target relevant to the token.

4a. The Perceived Moral Relationship

This section states the constitutive heart of the MDT. The downstream token rules are logically posterior.

4a.1 Frame-relative agent and object construals

A frame presents an agent referent and an object referent under particular construals. The same person, group, artifact, institution, ideal, or event may therefore receive different builder profiles under different frames. Agent and object construals are represented bindings, not metaphysical claims.

AgentConstrual_s(x,y_A,f,c) ObjectConstrual_s(o,y,f,c)

Role co-reference. NDT does not require the agent referent and object referent to be distinct. One thin referent may be bound under both a frame-relative agent construal and a frame-relative object construal:

AgentConstrual_s(x_self,y_self,f,c) and ObjectConstrual_s(o_self,y_self,f,c)

Where that role co-reference participates in a qualifying relationship and in the ordinary same-chain provenance conditions, the result is moral self-relation. Co-reference is not role identity: the agent construal presents the represented subject of action, address, or response, while the object construal presents the standing, condition, treatment, or interest around which the relationship is organized. A route that morally types the token must still carry role-correct subject and object bindings; one binding token does not automatically substitute for the other.

The same broad case and represented execution may support distinct judgment tokens under different object construals. A reconstruction that maps "I killed her" and "I killed myself" onto one execution may therefore share an execution bridge while changing which object organizes the appraisal. Whether the resulting pair is comparison-ready still depends on the shared question or dimension bridge required by §6.

The core perceived moral relationship has three builders. Inherent valence asks whether o is represented as having positive inherent value, negative inherent value, both, or neither. The positive and negative components form one profile but remain independently represented; they are not endpoints of one scale and may not be cancelled into a net score. Perceived agency asks whether y_A-under-x is represented as having at least one object- or relationship-relevant response affordance within the token's declared modal-temporal horizon. Relational Positioning asks how the exact represented object is located relative to the exact represented subject through an identity-mediated, intimacy, or direct-encounter bridge. It is valence-neutral, non-scalar, and independently variable from the inherent-valence profile. Third-personal subjecthood is gated separately by attributed moral vision. A bound relationship token and profile-sensitive actual contribution are further universal conditions, but they are not additional builders.

The builder argument is conditional and contrastive. Given NDT's explanandum, removing the inherent-valence profile erases the distinction between a represented object that matters in itself positively or negatively and a merely instrumental or valence-null object. Removing perceived agency erases the distinction between passive regard and the subject's situated capacity to affect or respond. Removing Relational Positioning removes every governed bridge that locates this exact object relative to this exact subject, leaving valence and agency co-present but not joined as a represented practical relationship. The deletion argument establishes conditional nonredundancy within NDT's present decomposition while the remaining construal, subject, token, and provenance machinery is held fixed. It does not establish that every adequate rival must use NDT's labels or variables. A rival may redistribute or combine the work if it preserves the distinctions and achieves better comparative adequacy.

Functional access features - including interface fit and availability to memory, attention, or awareness - help a representation enter a possible contributing process. They belong to access and actual-contribution machinery, not to Relational Positioning. Represented distance is the absence of a Relational Positioning witness under an adequate detection scheme, not a fourth route or a negative positioning value. Represented opposition is a contrastive-identity position and does not by itself entail negative inherent value. The former standing-mediated category is retired: an impersonal directed assignment qualifies, when it qualifies at all, through an operative complementary-role identity bridge; downstream duties or permissions supported by that identity are credited separately as bender edges.

- Positive inherent value PositiveInherentValue_s(o,t,f,c): the degree to which o is represented as mattering positively or as good in itself beyond merely instrumental use.

- Negative inherent value $\text{NegativeInherentValue}_s(o,t,f,c)$: the degree to which o is represented as mattering negatively or as bad in itself beyond merely instrumental disvalue.

Neither component entails any particular requirement, permission, priority, appraisal, blame attribution, avoidance, punishment, promotion, preservation, or other downstream response. Those consequences require separately represented and credited edges.

$\text{InherentValenceProfile}_s(o,t,f,c) := \langle \text{PositiveInherentValue}_s(o,t,f,c), \text{NegativeInherentValue}_s(o,t,f,c) \rangle$

$\text{PositiveOnlyValence}_s(o,t,f,c) \leftrightarrow \text{PositiveInherentValue}_s(o,t,f,c) > 0 \wedge \text{NegativeInherentValue}_s(o,t,f,c) = 0$

$\text{NegativeOnlyValence}_s(o,t,f,c) \leftrightarrow \text{PositiveInherentValue}_s(o,t,f,c) = 0 \wedge \text{NegativeInherentValue}_s(o,t,f,c) > 0$

$\text{MixedValence}_s(o,t,f,c) \leftrightarrow \text{PositiveInherentValue}_s(o,t,f,c) > 0 \wedge \text{NegativeInherentValue}_s(o,t,f,c) > 0$

$\text{NeitherValence}_s(o,t,f,c) \leftrightarrow \text{PositiveInherentValue}_s(o,t,f,c) = 0 \wedge \text{NegativeInherentValue}_s(o,t,f,c) = 0$

These four states are exhaustive and mutually exclusive for the bounded two-component profile. A governed empirical instance supplies prospective positive and negative branch predicates $\mathcal{U}_+$ and $\mathcal{U}_-$. Each branch receives only its corresponding component; neither branch receives a maximum, difference, sum, mean, or other scalar reduction of the pair. The positive branch must preserve the frozen 1.5.3 positive-only *final-value branch* verdict. This is not whole-core equivalence: it does not compare the old agency gate, relationship interaction, Relational Positioning predecessor, subject gate, or moral verdict. The negative branch must be independently specified and sensitive to negative inherent value in its declared range. Both branches are false at zero. A finite implementation may use replaceable component floors and combine their Boolean results only at the displayed decision boundary, but NDT fixes no universal positive or negative threshold.

$\text{FrozenLegacyFinalValueFloor153}_s(K,c)$: the prospectively registered final-value component floor in the exact 1.5.3 model instance for comparison class K and context c . This historical parameter is used only for the narrow compatibility check; it is not a universal threshold or a current builder.

$\text{FrozenLegacyPositiveOnlyProjection153}_s(o,t,f,c) := \text{PositiveInherentValue}_s(o,t,f,c)$ when $\text{NegativeInherentValue}_s(o,t,f,c) = 0$

$\text{FrozenLegacyPositiveValuePass153}_s(v,K,c) \leftrightarrow v \geq \text{FrozenLegacyFinalValueFloor153}_s(K,c)$

$\text{PositiveOnlyValenceBranchEquivalentTo153}_s(\mathcal{U}_{\text{plus}}, K, c) \leftrightarrow \forall v \in [0,1]$
 $[\mathcal{U}_{\text{plus}}(v) = \text{FrozenLegacyPositiveValuePass153}_s(v, K, c)]$

$\text{ValenceBranchesAdmissible}_s(\mathcal{U}_{\text{plus}}, \mathcal{U}_{\text{minus}}, K, c) \leftrightarrow \text{ProspectivelyRegistered}_s(\mathcal{U}_{\text{plus}}, K, c) \wedge$
 $\text{ProspectivelyRegistered}_s(\mathcal{U}_{\text{minus}}, K, c) \wedge \mathcal{U}_{\text{plus}}(0) = \text{false} \wedge \mathcal{U}_{\text{minus}}(0) = \text{false} \wedge$
 $\text{PositiveOnlyValenceBranchEquivalentTo153}_s(\mathcal{U}_{\text{plus}}, K, c) \wedge \text{NegativeBranchSensitive}_s(\mathcal{U}_{\text{minus}}, K, c) \wedge$
 $\text{ProfileCategoryPreserved}_s(\mathcal{U}_{\text{plus}}, \mathcal{U}_{\text{minus}}, K, c)$

$\text{PositiveValenceBranchPass}_s(o,t,f,c) \leftrightarrow \text{PositiveValenceCriterionOf}_s(o,t,f,c)(\text{PositiveInherentValue}_s(o,t,f,c)) = \text{true}$

$\text{NegativeValenceBranchPass}_s(o,t,f,c) \leftrightarrow \text{NegativeValenceCriterionOf}_s(o,t,f,c)(\text{NegativeInherentValue}_s(o,t,f,c)) = \text{true}$

$\text{ValenceGate}_s(o,t,f,K,c) \leftrightarrow \text{ValenceBranchesAdmissible}_s(\text{PositiveValenceCriterionOf}_s(o,t,f,c),$
 $\text{NegativeValenceCriterionOf}_s(o,t,f,c), K, c) \wedge (\text{PositiveValenceBranchPass}_s(o,t,f,c) \vee \text{NegativeValenceBranchPass}_s(o,t,f,c))$

- Represented agency-capacity $\text{AgencyCap}_s(x,o,t,f,c)$: the degree to which s attributes to y_A -under- x at least one object- or relationship-relevant response affordance within the declared modal-temporal horizon.

$\text{ModalTemporalHorizonOf}_s(x,o,t,f,c) = h$: the horizon fixed for the token before classification. It specifies which represented present, prospective, historical, counterfactual, collective, role-mediated, reparative, testimonial, or expressive response possibilities are eligible. The horizon is not expanded after observing the target verdict.

$\text{ResponseAffordance}_s(r_A, x, y_A, o, y, h, f, c)$: under construal x in frame f , s represents agent referent y_A as able, within h , to perform or sustain response r_A that bears on object referent y under construal o or on the bound relationship to that object.

$\text{AgencyCap}_s(x,o,t,f,c) := \text{Degree}_s[\{r_A : \text{exists } y[\text{ObjectConstrual}_s(o,y,f,c) \text{ and}$
 $\text{ResponseAffordance}_s(r_A, x, \text{AgentAnchorOf}_s(x,f,c), o, y, \text{ModalTemporalHorizonOf}_s(x,o,t,f,c), f, c)] \}]$

AgencyCap is representational. A falsely attributed capacity can satisfy the constitutive role even when the capacity is absent in fact. Present feasibility of a particular act is represented separately by Feas_delta ; awareness, evaluation, subjecthood, or the bare possibility of thought does not by itself supply a response affordance. If the active model represents total inability to affect, answer, repair, testify, express, or otherwise respond across every route admitted by the fixed horizon, AgencyCap is zero for that token. Impossible-duty cases at that boundary remain an explication-adequacy question rather than an automatic redefinition of the horizon.

$\text{FunctionalAccess}_s(u0, \psi, c)$: active unit $u0$ is available to process ψ under the declared access model in c . Functional access is necessary for the affected route to contribute, but it neither establishes concern nor confers moral type.

$\text{RelationalPositioning_s}(\lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \text{IdentityMediatedRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \vee$
 $\text{IntimacyRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \vee \text{DirectEncounterRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$

$\text{RelationalPositioningGate_s}(m, x, y_A, o, y, t, f, c) \leftrightarrow \exists \lambda \text{ RelationalPositioning_s}(\lambda, m, x, y_A, o, y, t, f, c)$

Identity-mediated Relational Positioning. Identity-mediated Relational Positioning is realized when an independently specified and operative identity or role schema locates the exact represented object relative to the exact represented subject and contributes through that bridge to the relationship's operative evaluation. Its governed forms are shared identity, contrastive identity, and complementary-role identity. Mere category membership, role possession, schema availability, or conventional association with duties is insufficient. Any requirement, permission, priority, appraisal, or other downstream consequence supported by the identity representation must be represented and credited separately as a strict-bender edge when the core is preserved, or as a core-changing edge otherwise. Identity-mediated positioning, inherent valence, and downstream bending remain independently variable.

$\text{IdentityBridgeQualified_s}(\iota, \lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \text{IndependentlySpecifiedIdentitySchema_s}(\iota, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{OperativeIdentitySchema_s}(\iota, m, f, c) \wedge \text{LocatesExactObjectRelativeToExactSubject_s}(\iota, x, y_A, o, y, f, c) \wedge$
 $\text{IdentityBridgeContributes_s}(\iota, \lambda, m, t, f, c)$

$\text{SharedIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \exists \iota [\text{IdentityBridgeQualified_s}(\iota, \lambda, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{SharedIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)]$

$\text{ContrastiveIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \exists \iota [\text{IdentityBridgeQualified_s}(\iota, \lambda, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{ContrastiveIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)]$

$\text{ComplementaryRoleIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \exists \iota [\text{IdentityBridgeQualified_s}(\iota, \lambda, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{ComplementaryRoleIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)]$

$\text{IdentityMediatedRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \leftrightarrow \text{SharedIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \vee$
 $\text{ContrastiveIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c) \vee \text{ComplementaryRoleIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$

The three identity forms are governed as follows. Shared identity locates the object and subject through an operative common identity, not bare co-membership. Contrastive identity locates them through an operative self/other, in-group/out-group, ally/opponent, or other identity-defined contrast; it has no fixed valence. Complementary-role identity locates them through mutually defining or directionally paired roles such as physician/patient, fiduciary/beneficiary, official/constituent, or teacher/student. The roles need not be consciously endorsed by both parties; the modeled subject must take the identity or role schema up as operative in the represented reasoning chain.

$\text{SharedIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)$: ι presents an operative common identity locating the exact object and subject; bare co-membership is insufficient.

$\text{ContrastiveIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)$: ι presents an operative identity-defined contrast locating the exact object and subject; the contrast has no fixed inherent valence.

$\text{ComplementaryRoleIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c)$: ι presents mutually defining or directionally paired roles locating the exact object and subject; possession of a role label without operative uptake is insufficient.

$\text{ImpersonalDirectedAssignment_s}(\iota, x, y_A, o, y, f, c) \wedge \text{IdentityBridgeQualified_s}(\iota, \lambda, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{ComplementaryRoleIdentityForm_s}(\iota, m, x, y_A, o, y, t, f, c) \rightarrow \text{ComplementaryRoleIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$

$\text{IntimacyRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$: route witness λ organizes the exact relationship through a particularized tie such as shared history, trust, dependence, relationship-specific knowledge, or mutual recognition.

$\text{DirectEncounterRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$: route witness λ organizes the exact relationship through the particular object's presence, condition, expression, testimony, or address becoming directly operative. Empathy may mediate encounter but is neither necessary nor sufficient.

Each of $\text{IdentityMediatedRP_s}$, IntimacyRP_s , and $\text{DirectEncounterRP_s}$ implies $\text{RelationalPositioning_s}$ for the same exact tuple. They are nonexclusive and individually optional. Their displayed disjunction is exhaustive for this version's governed route taxonomy; an extension requires a separately ratified successor specification and may not be introduced post hoc to rescue an adverse case.

$\text{AdequateRPScheme_s}(\mathbf{R}, K, c) \wedge \neg \exists \lambda \text{ RelationalPositioning_s}(\lambda, m, x, y_A, o, y, t, f, c) \rightarrow$
 $\text{RelationalPositioningAbsent_s}(m, x, y_A, o, y, t, f, c)$

$\text{ContrastiveIdentityRP_s}(\lambda, m, x, y_A, o, y, t, f, c)$ does not imply $\text{NegativeInherentValue_s}(o, t, f, c) > 0$

$\text{RelationalSubrouteOf}(\lambda, \text{chi}, c)$: λ is a relational-route subgraph realized inside $\text{ChainGraph}(\text{chi})$

$\text{RelationalPositioningOn_s}(\lambda, m, x, y_A, o, y, t, f, \text{chi}, c) \leftrightarrow \text{RelationalPositioning_s}(\lambda, m, x, y_A, o, y, t, f, c)$

$\wedge \text{RelationalSubrouteOf}(\lambda, \text{chi}, c)$

$\wedge m \in \text{Nodes}(\text{lambda})$

RelationalPositioning_s and RelationalPositioningGate_s state the constitutive umbrella role. Route indicators, identity-schema indicators, and the adequate negative-detection scheme are provisional implementation and identification proposals. Functional access, salience, familiarity, generic association, and category membership do not imply positioning.

A self-described Kantian or utilitarian judgment is moral only if a qualifying relationship token actually organizes its evaluation. When reasoning consists only in manipulating a rule, utility function, or institutional policy, NDT permits the output to be normative but nonmoral. This is a discriminating consequence, not a case the architecture must be altered to avoid.

4a.2 Moral vision and third-personal moral subjecthood

Moral vision must not be defined as "the capacity to see a moral object," because that would make the analysis circular. Instead, define a nonmoral integrative competence and distinguish the referent to which the capacity is attributed from the construal under which it is modeled. CapableAs_s(y_A,x,Φ,c) means that s attributes to agent referent y_A, under agent construal x, the disposition to instantiate profile Φ under suitable enabling conditions. IRAttributed_s(y_A,x,o,t,f,c') is the corresponding integrative-recognition profile. This attribution becomes necessary when a judgment treats a represented nonself agent as a moral subject; it is not required merely for the judging system to stand in a moral relationship with an object of concern.

CapableAs_s(y_A,x,Φ,c): s attributes to agent referent y_A, under construal x, the capacity to instantiate profile Φ under suitable conditions

IRAttributed_s(y_A,x,o,t,f,c') $\leftrightarrow$ RepresentsInherentValenceAs_s(y_A,x,o,t,f,c')

$\wedge$ RepresentsRelevantRelationsAs_s(y_A,x,o,t,f,c')

$\wedge$ RepresentsAgencyModelAs_s(y_A,x,o,t,f,c') $\wedge$ IntegratesToEvalStructureAs_s(y_A,x,o,t,f,c')

MV_s(x,o,t,f,c) := Degree_s[CapableAs_s(AgentAnchorOf_s(x,f,c),x, λc'.IRAttributed_s(AgentAnchorOf_s(x,f,c),x,o,t,f,c'), c)]

No moral predicate occurs on the right-hand side. AgentAnchorOf is well defined by the functional AgentConstrual relation above. CapableAs attributes a capacity to that agent referent as presented under x; it does not treat the construal token itself as a bearer of psychological capacities. The frame index identifies the agent and object construals relative to which s makes the attribution, and it does not require the underlying agent referent consciously to represent that frame label. The label moral vision applies because this otherwise nonmoral competence gates third-personal moral subjecthood.

4a.3 Builder gates, subject gate, and governed empirical instances

Normalize PositiveInherentValue, NegativeInherentValue, and AgencyCap to [0,1] inside a declared empirical instance. The constitutive source does not fix a universal positive or negative valence threshold. It instead constrains the prospectively registered positive and negative branch predicates stated in §4a.1. Perceived agency retains its own noncompensable floor. Neither branch nor the agency gate accepts a scalar reduction of the inherent-valence profile. Relational Positioning has its own existential gate requiring at least one bound positioning witness.

CoreGate_s(x,o,t,f,K,c) $\leftrightarrow$ ValenceGate_s(o,t,f,K,c)

$\wedge$ AgencyCap_s(x,o,t,f,c) $\geq \theta_A$

The relationship interaction retained from 1.5.3 is represented here as an ordered two-branch profile rather than one aggregate. Each branch combines only its corresponding inherent-valence component with perceived agency. The positive and negative interaction functions and their branch criteria are prospectively registered for the empirical instance. The ordered pair remains available after the branch decision; only the two Boolean branch results may be disjoined at the displayed RelationshipStrengthGate_s boundary. No maximum, difference, sum, mean, or other scalar reduction of the pair is admissible.

PositiveRelationshipStrength_s(x,o,t,f,c) := PositiveAgencyInteractionOf_s(x,o,t,f,c)
(PositiveInherentValue_s(o,t,f,c), AgencyCap_s(x,o,t,f,c))

NegativeRelationshipStrength_s(x,o,t,f,c) := NegativeAgencyInteractionOf_s(x,o,t,f,c)
(NegativeInherentValue_s(o,t,f,c), AgencyCap_s(x,o,t,f,c))

RelationshipStrengthProfile_s(x,o,t,f,c) := <PositiveRelationshipStrength_s(x,o,t,f,c),
NegativeRelationshipStrength_s(x,o,t,f,c)>

FrozenLegacyRelationshipInteraction153_s(v,a,K,c): the value returned for inputs v and a by the prospectively registered 1.5.3 relationship-interaction function in comparison class K and context c.

FrozenLegacyRelationshipStrengthBoundary153_s(t,f,K,c): the prospectively registered 1.5.3 relationship-strength boundary for time t and frame f in comparison class K and context c.

FrozenLegacyPositiveRelationshipPass153_s(v,a,t,f,K,c) $\leftrightarrow$ FrozenLegacyRelationshipInteraction153_s(v,a,K,c) $\geq$ FrozenLegacyRelationshipStrengthBoundary153_s(t,f,K,c)

PositiveOnlyRelationshipBranchEquivalentTo153_s(x,o,t,f,K,c) $\leftrightarrow \forall v,a \in [0,1]$
 [PositiveRelationshipStrengthCriterionOf_s(x,o,t,f,c)(PositiveAgencyInteractionOf_s(x,o,t,f,c)(v,a))=
 FrozenLegacyPositiveRelationshipPass153_s(v,a,t,f,K,c)]

RelationshipStrengthProfilePreserved_s(x,o,t,f,K,c) $\leftrightarrow$ ProspectivelyRegistered_s(RelationshipStrengthProfile_s(x,o,t,f,c),K,c)

RelationshipStrengthBranchesAdmissible_s(x,o,t,f,K,c) $\leftrightarrow$
 ProspectivelyRegistered_s(PositiveAgencyInteractionOf_s(x,o,t,f,c),K,c) $\wedge$
 ProspectivelyRegistered_s(NegativeAgencyInteractionOf_s(x,o,t,f,c),K,c) $\wedge$
 ProspectivelyRegistered_s(PositiveRelationshipStrengthCriterionOf_s(x,o,t,f,c),K,c) $\wedge$
 ProspectivelyRegistered_s(NegativeRelationshipStrengthCriterionOf_s(x,o,t,f,c),K,c) $\wedge$
 PositiveOnlyRelationshipBranchEquivalentTo153_s(x,o,t,f,K,c) $\wedge$ NegativeRelationshipBranchSensitive_s(x,o,t,f,K,c) $\wedge$
 RelationshipStrengthProfilePreserved_s(x,o,t,f,K,c)

PositiveRelationshipStrengthBranchPass_s(x,o,t,f,c) $\leftrightarrow$ PositiveRelationshipStrengthCriterionOf_s(x,o,t,f,c)
 (PositiveRelationshipStrength_s(x,o,t,f,c))=true

NegativeRelationshipStrengthBranchPass_s(x,o,t,f,c) $\leftrightarrow$ NegativeRelationshipStrengthCriterionOf_s(x,o,t,f,c)
 (NegativeRelationshipStrength_s(x,o,t,f,c))=true

RelationshipStrengthGate_s(x,o,t,f,K,c) $\leftrightarrow$ RelationshipStrengthBranchesAdmissible_s(x,o,t,f,K,c) $\wedge$
 (PositiveRelationshipStrengthBranchPass_s(x,o,t,f,c) $\vee$ NegativeRelationshipStrengthBranchPass_s(x,o,t,f,c))

CoreMRCondition_s(x,o,t,f,K,c) $\leftrightarrow$ CoreGate_s(x,o,t,f,K,c)

$\wedge$ RelationshipStrengthGate_s(x,o,t,f,K,c)

BuilderQualifiedMRCondition_s(m,x,y_A,o,y,t,f,K,c) $\leftrightarrow$ CoreMRCondition_s(x,o,t,f,K,c)

$\wedge$ RelationalPositioningGate_s(m,x,y_A,o,y,t,f,c)

The relationship-strength profile is a derived interaction over two existing builders, inherent valence and perceived agency; it is not a fourth builder. Its positive-only compatibility clause compares only the former 1.5.3 relationship-strength branch after the old component gates. It does not claim whole-core equivalence and does not import the old Closeness, subject, token, or moral-verdict conditions. No universal interaction function or strength threshold is fixed for current negative-only or mixed profiles.

The three builder requirements are conjunctive but not definitionally interchangeable. For a fixed well-typed relationship tuple, admissible finite assignments can make CoreMRCondition true and RelationalPositioningGate false, RelationalPositioningGate true and CoreMRCondition false, or both true. Functional access is separately route- and process-indexed; neither gate entails AccessFor_s. These are formal independence claims inside the explication, not empirical prevalence claims.

One constructive assignment uses $\theta A=0.5$ and admissible finite valence and relationship-strength branch criteria whose decision boundaries pass the positive-only profile $\langle 1,0 \rangle$ with AgencyCap=1 and fail $\langle 0,0 \rangle$. With $\langle 1,0 \rangle$, AgencyCap=1, both derived gates passing, and no RelationalPositioning witness, CoreMRCondition is true while RelationalPositioningGate is false. With $\langle 0,0 \rangle$, AgencyCap=1, and one ComplementaryRoleIdentityRP witness, RelationalPositioningGate is true while CoreMRCondition is false. With $\langle 1,0 \rangle$, AgencyCap=1, and that witness, both hold. A fourth assignment holds the positive inherent-valence component and AgencyCap fixed while varying the negative inherent-valence component across the registered sensitivity witnesses of both negative branches. No ordinary-language case verdict establishes these independence results.

CoreMRCondition is a satisfaction condition in the theorist's metalanguage. It is not itself a premise token. The valence gate prevents a neither profile from passing; the agency floor prevents extreme valence from compensating for absent perceived agency; the relationship-strength gate requires a registered valence-specific agency interaction; and noncancellation prevents either branch from erasing the other. No quantitative setting can compensate for absent Relational Positioning. Third-personal moral subjecthood is governed separately:

MoralSubjectGate_s(x,y_A,o,t,f,c) $\leftrightarrow$ SelfReferent_s(y_A,c)

$\vee MV_s(x,o,t,f,c) \geq \theta M$

The first disjunct permits first-personal moral judgment without a separately attributed self-directed moral-vision premise. The second preserves the shark/competent-human contrast for third-personal moral appraisal.

Self-reference clears only the subject gate. When $\text{SelfReferent}_s(y_A, c)$ holds, NDT does not require a separately attributed self-directed moral-vision premise. Self-reference does not by itself imply **CoreGate**, **RelationalPositioningGate**, **BuilderQualifiedMRCondition**, **Outputs**, or **MoralJudgment**. A self-directed moral token must still clear the valence decision boundary and agency gate, instantiate some **Relational Positioning** realization, contain a represented relationship token with role-correct self subject and self object bindings, and satisfy the ordinary same-chain provenance conditions.

This permits moral concern for oneself without making every self-directed prudential or normative judgment moral. It also permits self-alienation or low self-directed visibility when the object-side profile fails.

4a.4 Represented relationship tokens and empirically identified instances

Let $\text{RelationshipRep}_s(m, x, y_A, o, y, t, f, c)$ mean that m is an active token in s 's representational economy presenting agent construal x of agent referent y_A as standing toward object construal o of object referent y at temporal anchor t under f . The token may be implicit and subpersonal; it need not be a consciously entertained proposition.

$\text{MoralRelationRep}_s(m, x, y_A, o, y, t, f, K, c) \leftrightarrow \text{RelationshipRep}_s(m, x, y_A, o, y, t, f, c)$

$\wedge \text{AgentConstrual}_s(x, y_A, f, c) \wedge \text{ObjectConstrual}_s(o, y, f, c)$

$\wedge \text{BuilderQualifiedMRCondition}_s(m, x, y_A, o, y, t, f, K, c)$

$\wedge \text{MoralSubjectGate}_s(x, y_A, o, t, f, c)$

$\text{MR}_s(x, y_A, o, y, t, f, K, c) \leftrightarrow \exists m \text{MoralRelationRep}_s(m, x, y_A, o, y, t, f, K, c)$ [extensional shorthand]

Only representation token m , not formula **CoreMRCondition**, **RelationalPositioningGate**, **BuilderQualifiedMRCondition**, **MoralSubjectGate**, or shorthand MR_s , can occur in $\text{Premises}(\rho)$. Because the satisfaction conditions are indexed to f and include both construal relations, a relationship token in one frame does not automatically confer moral status on a token involving another construal of the same agent or object referent.

The governed valence branches, relationship-interaction branches, and separate agency and subject-gate thresholds are constitutive model schemas rather than point-identified quantitative models. Particular empirical instances must fix:

$\mathcal{M}_R := \langle \mathcal{V}_{\text{plus}}, \mathcal{V}_{\text{minus}}, \theta_A, \mathcal{J}_{\text{plus}}, \mathcal{J}_{\text{minus}}, \mathcal{G}_{\text{plus}}, \mathcal{G}_{\text{minus}}, \mathcal{R} \rangle$ $\mathcal{M}_S := \langle \theta_M \rangle$

Here $\mathcal{J}_{\text{plus}}$ and $\mathcal{J}_{\text{minus}}$ are the positive-by-agency and negative-by-agency interaction functions; $\mathcal{G}_{\text{plus}}$ and $\mathcal{G}_{\text{minus}}$ are their separate Boolean branch criteria. Each item is prospectively registered.

An empirical application must also pre-specify an admissible identification model $\mathcal{R}$ for the non-scalar **RelationalPositioning** predicate and its exact agent, object, referent, time, frame, and relationship-token bindings. Positioning may not be inferred from the target's moral label, valence, a downstream duty, or the presence of an untested category or role.

For quantitative testing, $\mathcal{V}_{\text{plus}}$, $\mathcal{V}_{\text{minus}}$, $\mathcal{J}_{\text{plus}}$, $\mathcal{J}_{\text{minus}}$, $\mathcal{G}_{\text{plus}}$, and $\mathcal{G}_{\text{minus}}$ should belong to pre-specified bounded-complexity families; the branch criteria and agency minimum should be calibrated on separate data or independently constrained; the moral-vision minimum should be calibrated independently for third-personal subjecthood; and frame or context variation should follow a pre-specified hierarchical rule. Parameters may not be adjusted case by case after observing the relationship or subjecthood signature.

Without fixed instances $\mathcal{M}_R$ and $\mathcal{M}_S$, bounded pre-specified families of instances, and a pre-specified **Relational Positioning** identification model, the schema is not operationally identified and should not be treated as yielding quantitative predictions or token classifications. Frame-relative inputs and multiple realization must not become excuses for post hoc parameter or route fitting.

Branch-preservation adequacy. Every finite instance must retain the ordered positive/negative inherent-valence pair and the ordered relationship-strength pair $\langle r_{\text{plus}}, r_{\text{minus}} \rangle$ through measurement and profile-sensitive contribution tests. It may compute positive-branch and negative-branch Boolean results for each profile and disjoin them only at **ValenceGate** and **RelationshipStrengthGate**, respectively. It may not substitute a maximum, net difference, sum, mean, or another one-scalar summary for either pair in any builder gate or relationship aggregation. Mixed, positive-only, negative-only, and neither must remain distinguishable after both gate verdicts.

4a.5 Access, route uptake, component diagnostics, and subjecthood

$\text{CoreRelationDiagnosticProfile}_s(x,o,t,f,c) := \langle \text{PositiveInherentValue}_s(o,t,f,c), \text{NegativeInherentValue}_s(o,t,f,c), \text{AgencyCap}_s(x,o,t,f,c), \text{RelationshipStrengthProfile}_s(x,o,t,f,c) \rangle$

Access, Relational Positioning uptake, route-family and identity-form uptake, and the component-preserving $\text{CoreRelationDiagnosticProfile}$ may be useful empirical diagnostics, but none is assumed to be a sufficient statistic. Two frame-relative cases with matched diagnostic profiles may differ in access, positioning organization, or actual contribution. Moral vision remains a separate diagnostic of third-personal moral subjecthood rather than a component hidden inside the component profile. No one-scalar relation-strength, distance, or closeness quantity is part of the current constitutive core. Relationship strength is an ordered, branch-preserving derived interaction profile; Relational Positioning is the non-scalar existence and organization condition stated by $\text{RelationalPositioningGate}$.

4a.6 The shark, the friend, the stranger, and reframed agents and objects

Hold the inherent-valence profile, perceived agency, the bound Relational Positioning realization, requirement, and support graph fixed. Under one frame, the judging system attributes insufficient integrative-recognition capacity to the represented nonself subject. The subject gate fails. The example referent may be a shark, but shark-kind does no formal work; the result is indexed to the referent-under-construal in that frame.

Under another frame, a human or nonhuman referent may be represented as capable of recognizing and integrating the object's standing, and the subject gate may clear. Conversely, a human referent represented as lacking that capacity fails. The intended contrast is high versus low attributed recognition under otherwise matched structure, not human versus nonhuman entity kinds. Answerability and requirement-specific responsibility remain additional, independently represented conditions for blame. In the competent-stranger example, the exact stranger-child three-builder profile may remain fixed while the observer represents the stranger as practically capable, recognition-capable, and answerable to a no-harm requirement.

$m_stranger\text{-}child \rightarrow u_no\text{-}harm \rightarrow \text{Acc}_{\delta S} \rightarrow \text{Wrong}(x_stranger, a_attack, t) \rightarrow \text{Blameworthy}(x_stranger, a_attack, t)$, where answerability and responsibility hold

The moral witness is the observer's represented stranger-child relationship, not merely the observer's parent-child relationship. The latter may intensify salience or special protective obligation; the former grounds what the observer represents the stranger as owing to the child. This is a substantive relational commitment: third-personal moral judgment represents the judged agent as standing under a moral relation to the relevant object, whether or not that agent endorses or consciously recognizes it.

Failures of access, Relational Positioning, or relationship-profile contribution are equally substantive. If an object's represented standing cannot enter any live route, if the exact relationship lacks a positioning realization, or if the relationship enters only as background information that does not help form the evaluation, the resulting token does not acquire moral type through that relationship. NDT treats such failures as possible structures of moral exclusion rather than as technical anomalies. The empirical burden is to identify access, positioning route and identity-form uptake, the three-builder profile, and actual contribution independently rather than inferring them from a verdict.

4a.7 Constitutive claim and scope

Claim-status decomposition. CE1 is the constitutive explication of the three builder roles. A2 and A3 are explicit assumptions proposing the represented subject/object necessities and the representationally agentless exclusion; they are not established empirical findings and do not equate moral subjecthood with ordinary agency. The full MDT typing clause is the model-relative constitutive rule conditional on CE1, A2, A3, and the remaining displayed provenance conditions. Its guarantees are derived only inside that architecture; empirical kind adequacy is separate and open.

Relative to a fixed model instance and host frame, an active relationship token represents a moral relationship exactly when its agent and object construals are bound, the governed inherent-valence branch criteria and agency floor are satisfied, at least one exactly bound positioning realization satisfies $\text{RelationalPositioningGate}$, and the self/nonself subject gate is satisfied. The self disjunct permits first-personal judgment without a redundant self-attribution of moral vision; the nonself disjunct requires attributed moral vision.

NDT uses moral cognition as an explanatory construct, not as a report of every ordinary, disciplinary, or philosophical use of the word moral. Family-resemblance variation in usage is compatible with a more precise cognitive kind. A label such as deontological, consequentialist, caring, sacred, or moral therefore does not decide whether each token satisfies the model's conditions.

The definitions are formally non-self-referential, and NDT proposes on that basis a conceptually noncircular distinction. That does not establish validated, operationally independent identification: target leakage, target-derived paradigms, and

post-hoc implicit subjects remain open methodological risks. Future blind, held-out, and target-leakage checks are examples for the identification program, not present evidence.

Kind adequacy is assessed by noncircular specification, coverage of paradigm and contrast cases, discrimination between same-content tokens with different provenance, integration with general normativity, comparative explanatory power, and consequences specified independently of the desired classification. Constitutive status does not immunize an explication from assessment. Sustained failure on independently specified contrast or consequence tests can require restriction, revision, or abandonment of the proposed construct. Where adequate measurement is not yet available, that is a registered limitation rather than evidence in the theory's favor. A contested case need not arrive with a fully articulated superior rival before it can become adverse evidence, but no single lexical verdict is automatically decisive. Section 8.4 states the direct represented-subject pressure test.

Concrete individuals are not privileged. A relationship may involve an abstract, collective, future, institutional, or ideal object, but no label such as humanity, society, rational nature, welfare, or the greater good automatically supplies a relationship. The token, object construal, three-builder profile, and provenance must be independently instantiated. Conversely, if a self-described Kantian, consequentialist, contractualist, or religious reasoner manipulates principles without such provenance, NDT classifies the reasoning as normative but nonmoral.

The account is descriptive and subject-relative. The present core formalism represents positive inherent value and negative inherent value as independent components of one profile. Mixed, positive-only, negative-only, and neither remain distinguishable throughout; otherness is separately represented by contrastive identity and has no fixed valence. Specific judgments additionally require a formed decision model, a typed candidate, and at least one realized live chain in the token's actual support graph.

4b. From Moral Relationship to Moral Judgment

Section 4a supplies the relationship schema. Section 4b states the additional provenance and subject-matching conditions under which a normative token receives moral type.

4b.1 Connected provenance, object bearing, and represented moral subjects

$\text{RequirementOf}(\kappa) = u \leftrightarrow \text{exists } x, a, t [\text{Content}(\kappa) = \text{Violation}(x, a, u, t)]$

$\text{AppraisalDimensionOf}(\kappa) = q \leftrightarrow [\text{exists } x, a, t [\text{Content}(\kappa) = \text{Praiseworthy}(x, a, t) \text{ and } \text{LocalDimensionOf}(\kappa) = q]] \text{ or } [\text{exists } z_A, \text{pol}, t [\text{Content}(\kappa) = \text{Appraise}(z_A, q, \text{pol}, t)]]$

$\text{RelevantEval}(\kappa) := \text{GuideStruct}(\text{BaseOf}(\kappa))$ when exists x, C, t $[\text{Content}(\kappa) = \text{OughtChoose}(x, C, t)]$;
 $\text{RequirementStruct}(\text{RequirementOf}(\kappa), \text{BaseOf}(\kappa))$ when exists x, a, u, t $[\text{Content}(\kappa) = \text{Violation}(x, a, u, t)]$;
 $\text{AcceptabilityStruct}(\text{BaseOf}(\kappa))$ when exists x, a, t $[\text{Content}(\kappa) = \text{Wrong}(x, a, t) \text{ or } \text{Content}(\kappa) = \text{Blameworthy}(x, a, t)]$;
 $\text{AppStruct_BaseOf}(\kappa)(\text{AppraisalDimensionOf}(\kappa))$ when exists x, a, t $[\text{Content}(\kappa) = \text{Praiseworthy}(x, a, t)]$ or exists z_A, q, pol, t $[\text{Content}(\kappa) = \text{Appraise}(z_A, q, \text{pol}, t)]$

The mutually exclusive judgment grammar and the typed Candidate clauses make RequirementOf , $\text{AppraisalDimensionOf}$, and RelevantEval functional on every output token. They are not asserted as total outside that domain.

$\text{EvaluationRelevant}(\phi, E, \kappa, c) := \phi$ is an independently specified realization-profile dimension whose value constitutes or functionally determines E for token κ in the declared model instance

$\text{EvalRelevantProfile_S}(E, \kappa, K, c) := \{ \langle \phi, v \rangle \text{ in } \text{RealizationProfile_S}(\kappa, K, c) : \text{EvaluationRelevant}(\phi, E, \kappa, c) \}$

For every output token κ , $S^* \{s, K^* \{ \kappa, c \}, c\}$ and $K^* _ \{ \kappa, c \}$ denote the registered constitutively adequate scheme and declared comparison class fixed prospectively for that declared class before any individual contribution verdict:

$\text{FixedConstitutiveScheme}(S^* \{s, K^* \{ \kappa, c \}, c\}, K^* \{ \kappa, c \}, \kappa, c) \leftrightarrow \text{FixedConstitutiveSchemeOf_s}(K^* \{ \kappa, c \}, c) = S^* \{s, K^* \{ \kappa, c \}, c\}$

$\text{IntervenesOnUnit}(I, u_0, \kappa, c)$: registered intervention I varies active support unit u_0 for token κ in c under the declared unit-intervention model

$\text{EvalProfileDifferenceMaker_}\{S^* \{s, K^* \{ \kappa, c \}, c\}\}$ $(u_0, E, \kappa, K^* \{ \kappa, c \}, c) \leftrightarrow \text{exists } I [I \text{ in } \text{RegisteredInterventions}(K^* \{ \kappa, c \}, c) \text{ and } \text{IntervenesOnUnit}(I, u_0, \kappa, c) \text{ and } \text{InvarianceConditionsHold}(I, K^* \{ \kappa, c \}, c) \text{ and } \text{EvalRelevantProfile}\{S^* \{s, K^* \{ \kappa, c \}, c\}\}(E, \kappa, I, K^* \{ \kappa, c \}, c) \neq \text{EvalRelevantProfile}\{S^* \{s, K^* \{ \kappa, c \}, c\}\}(E, \kappa, K^* _ \{ \kappa, c \}, c)]$

$\text{ProcessCarriesOn}(\psi, u_0, \chi, \kappa, c)$: active process ψ realizes the segment of live chain χ by which support unit u_0 can affect token κ

$\text{AccessFor_s}(u_0, \text{chi}, \text{kappa}, c) \leftrightarrow \text{exists } \text{psi} [\text{ProcessCarriesOn}(\text{psi}, u_0, \text{chi}, \text{kappa}, c) \text{ and } \text{FunctionalAccess_s}(u_0, \text{psi}, c)]$

$\text{LocalContributionVerdict_}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (u_0, E, \text{chi}, \text{kappa}, K^* \{ \text{kappa}, c \}, c) \leftrightarrow u_0 \text{ in } \text{Nodes}(\text{chi}) \text{ and } \text{Reach_chi}(u_0, E, c) \text{ and } \text{AccessFor_s}(u_0, \text{chi}, \text{kappa}, c) \text{ and } \text{MemberOfActualMinimalSufficientSet}(u_0, \text{EvalRelevantProfile}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (E, \text{kappa}, K^* \{ \text{kappa}, c \}, c) \text{ and } \text{EvalProfileDifferenceMaker}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (u_0, E, \text{kappa}, K^* \{ \text{kappa}, c \}, c))$

$\text{LocalActualContribution}(u_0, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{LocalContributionVerdict_}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (u_0, E, \text{chi}, \text{kappa}, K^* \{ \text{kappa}, c \}, c)$

$\text{EvalContributes}(m, E, \text{chi}, \text{kappa}, c) \leftrightarrow m \text{ in } \text{Premises}(\text{chi}) \text{ and } \text{LocalActualContribution}(m, E, \text{chi}, \text{kappa}, c)$

$\text{ObjectBindingNode_s}(n_o, o, y, \text{chi}, \text{delta}, c) \leftrightarrow n_o = \text{ObjectBindingClaim_s}(\text{delta}, o, y, c) \text{ and } n_o \text{ in } \text{Premises}(\text{chi})$

$\text{ObjectContributes_s}(o, y, E, \text{chi}, \text{kappa}, \text{delta}, c) \leftrightarrow \text{exists } n_o [\text{ObjectBindingNode_s}(n_o, o, y, \text{chi}, \text{delta}, c) \text{ and } \text{LocalActualContribution}(n_o, E, \text{chi}, \text{kappa}, c)]$

$\text{OrganizesOn}(m, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{RealizedLiveChain}(\text{chi}, \text{kappa}, c) \text{ and } \text{EvalContributes}(m, E, \text{chi}, \text{kappa}, c)$

$\text{BoundRelationalPositioning_s}(m, x, y_A, o, y, t, f, c) := \{ \text{lambda} : \text{RelationalPositioning_s}(\text{lambda}, m, x, y_A, o, y, t, f, c) \}$

$\text{BoundBuilderProfile_s}(m, x, y_A, o, y, t, f, c) := < \text{InherentValenceProfile_s}(o, t, f, c), \text{AgencyCap_s}(x, o, t, f, c), \text{BoundRelationalPositioning_s}(m, x, y_A, o, y, t, f, c) > \text{ when } \text{RelationshipRep_s}(m, x, y_A, o, y, t, f, c)$

$\text{ChangesBoundBuilderProfile_s}(I, m, x, y_A, o, y, t, f, c)$: I changes at least one inherent-valence component, perceived agency, or the non-scalar positioning organization component of that exact $\text{BoundBuilderProfile_s}$ tuple under the declared profile-intervention model

$\text{RelationshipBackgroundPreserved_s}(I, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c)$: I preserves that exact agent, agent referent, object, object referent, time, and frame binding; salience state; judgment-event anchor; case; the identity of chain chi outside the selected profile-bearing subroute; all route metadata not selected as the intervention target; and every other background feature registered for this comparison

$\text{NonProfileRelationshipFeaturesPreserved_s}(I, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c)$: every feature of m and its subroute on chi outside the exact $\text{BoundBuilderProfile_s}$ tuple remains fixed under I

$\text{NonProfileEvalDriversPreservedOn_s}(I, m, E, \text{chi}, \text{kappa}, c)$: every registered driver of E on chi that is not constituted by the exact bound builder profile of m remains fixed under I . In particular, a feature d that changes alongside the builder profile but is not part of that profile invalidates this clause.

$\text{InterventionCounterpartChain_s}(I, \text{chi}, \text{chi_I}, \text{kappa}, c)$: chi_I is the intervention counterpart of chi for $\text{kappa}|I$, preserving the chain's event, target, bindings, and nonintervened support structure. A chain unrelated to chi cannot satisfy this predicate.

$\text{BuilderProfileInterventionOn_s}(I, m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \leftrightarrow I \text{ in } \text{RegisteredInterventions}(K^* \{ \text{kappa}, c \}, c) \text{ and } \text{IntervenesOnUnit}(I, m, \text{kappa}, c) \text{ and } \text{ChangesBoundBuilderProfile_s}(I, m, x, y_A, o, y, t, f, c) \text{ and } \text{RelationshipBackgroundPreserved_s}(I, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c) \text{ and } \text{NonProfileRelationshipFeaturesPreserved_s}(I, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c) \text{ and } \text{NonProfileEvalDriversPreservedOn_s}(I, m, E, \text{chi}, \text{kappa}, c) \text{ and } \text{InvarianceConditionsHold}(I, K^* \{ \text{kappa}, c \}, c)$

$\text{EvalStateOn_}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (E, \text{chi}, \text{kappa}, K^* \{ \text{kappa}, c \}, c)$: the evaluation-relevant state of E realized through chi for kappa under the fixed constitutive scheme

$\text{ControlledSameChainProfileEffect_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{exists } I, \text{chi_I} [\text{BuilderProfileInterventionOn_s}(I, m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \text{ and } \text{InterventionCounterpartChain_s}(I, \text{chi}, \text{chi_I}, \text{kappa}, c) \text{ and } \text{EvalStateOn_}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (E, \text{chi_I}, \text{kappa}|I, K^* \{ \text{kappa}, c \}, c) \neq \text{EvalStateOn}\{S^* \{s, K^* \{ \text{kappa}, c \}, c \} \} (E, \text{chi}, \text{kappa}, K^* \{ \text{kappa}, c \}, c)]$

$\text{BoundBuilderProfileUnit_s}(p_m, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c)$: p_m is the exact $\text{BoundBuilderProfile_s}$ tuple of m as realized on the profile-bearing subroute of chi for kappa

$\text{ActualSufficientForOn}(A, E, \text{chi}, \text{kappa}, c)$: active support set A , restricted to chi , is sufficient for the evaluation-relevant state of E in the actual model instance for kappa

$\text{ProfileNESSOn_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{exists } p_m, A [\text{BoundBuilderProfileUnit_s}(p_m, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c) \text{ and } p_m \text{ in } A \text{ and } A \text{ subseq } \text{Nodes}(\text{chi}) \text{ and } \text{ActualSufficientForOn}(A, E, \text{chi}, \text{kappa}, c) \text{ and not } \text{ActualSufficientForOn}(A \setminus \{p_m\}, E, \text{chi}, \text{kappa}, c)]$

$\text{RelationshipProfileDifferenceMaker_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{ControlledSameChainProfileEffect_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c) \text{ or } \text{ProfileNESSOn_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c)$

$\text{RelationshipProfileContributes_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, K) \leftrightarrow \text{exists } \text{lambda} [\text{RelationalPositioningOn_s}(\text{lambda}, m, x, y_A, o, y, t, f, \text{chi}, c) \text{ and } \text{MoralRelationRep_s}(m, x, y_A, o, y, t, f, K, c) \text{ and } K = \text{ComparisonClassOf}(\text{kappa}, c) \text{ and } f = \text{FrameOf}(\text{BaseOf}(\text{kappa})) \text{ and } \text{OrganizesOn}(m, E, \text{chi}, \text{kappa}, c) \text{ and } \text{RelationshipProfileDifferenceMaker_s}(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, c)]$

This condition prevents separate witnesses from doing the required jobs. The inherent-valence profile, agency, and a nonempty positioning organization cannot qualify one relationship while an unrelated feature d that changed in the same intervention is the only feature that altered E . Nor can relationship token m contribute on chi_1 while the builder-profile effect occurs only on unrelated chi_2 . The controlled branch fixes nonprofile relationship features and nonprofile evaluative drivers and compares chi only with its registered intervention counterpart. The NESS-style branch preserves redundant support: the exact profile may be necessary within one actual sufficient support set even when another sufficient set also remains active. It therefore does not require every builder, positioning route or identity form, or redundantly supporting realization to be an individual but-for cause. The exact, object-bound three-builder profile belongs to relationship token m , a bound positioning-realization subroute must occur on chi , and m must locally actually contribute there. These are two distinct evidential levels. OrganizesOn requires the relationship token's tier-one LocalActualContribution, including its own registered local difference-maker verdict. ProfileNESSOn is an alternative only at tier two, where contribution is attributed to the exact bound builder profile despite redundant profile support. A true ProfileNESSOn predicate cannot rescue a case in which the relationship token's tier-one local difference-maker is false. Where neither a controlled same-chain effect nor a profile-level minimal-sufficient witness is presently identifiable, the provenance verdict is empirically unresolved rather than positive by default. Functional access remains route- and process-indexed through AccessFor_s.

BearsOnOn($kappa,o,y,E,chi,c$) $\leftrightarrow$ RealizedLiveChain($chi,kappa,c$) and
CarriesObjectBinding_s($ModelRoute(chi),BaseOf(kappa),o,y,c$) and ObjectContributes_s($o,y,E,chi,kappa,BaseOf(kappa),c$)

SubjectBindingNode_s($n,x,y_A,kappa,chi,c$): n is the exact typed subject-binding node on chi that presents agent referent y_A under construal x as the subject of the target appraised in $kappa$

SubjectContributes_s($x,y_A,E,chi,kappa,c$) $\leftrightarrow$ exists n [SubjectBindingNode_s($n,x,y_A,kappa,chi,c$) and
LocalActualContribution($n,E,chi,kappa,c$)]

IntegrationRoleFor($n_i,kappa,c$): n_i is an independently typed requirement, appraisal, acceptability, or other relation-bearing evaluation-input node for $kappa$; a bare guide, candidate, output, or post-hoc common descendant does not satisfy this role

BoundProfileIntegrationOn_s($n_i,m,n_o,E,chi,kappa,c$) $\leftrightarrow$ IntegrationRoleFor($n_i,kappa,c$) and $\{m,n,n_o,n_i,E\}$ subseq
Nodes(chi) and Reach_chi(m,n_i,c) and Reach_chi(n,n_i,c) and Reach_chi(n_o,n_i,c) and Reach_chi(n_i,E,c)

This clause rejects a loosely coupled case in which the relationship, subject binding, and object binding travel on parallel same-chain subroutes and meet only at a generic guide, candidate, or output. It does not require either binding node to be an ancestor of the relationship node; convergence at one independently typed relation-bearing evaluation input is sufficient.

ExplicitMoralSubjectOn_s($kappa,x,y_A,chi,c$) $\leftrightarrow$ AgentIn(Content($kappa$), x) and AgentReferentOf(BaseOf($kappa$))= y_A and
CarriesAgentBinding_s($ModelRoute(chi),BaseOf(kappa),y_A,c$) and
SubjectContributes_s($x,y_A,RelevantEval(kappa),chi,kappa,c$)

ImplicitMoralSubjectOn_s($kappa,x,y_A,chi,c$) $\leftrightarrow$ Content($kappa$)=Appraise(z_A,q,pol,t) and not exists x'
AgentIn(Content($kappa$), x') and CarriesSubjectBinding_s($chi,kappa,x,y_A,c$) and SubjectTargetsOn_s(x,y_A,z_A,chi,c) and
SubjectContributes_s($x,y_A,RelevantEval(kappa),chi,kappa,c$)

MoralSubjectOf_s($kappa,x,y_A,chi,c$) $\leftrightarrow$ ExplicitMoralSubjectOn_s($kappa,x,y_A,chi,c$) or
ImplicitMoralSubjectOn_s($kappa,x,y_A,chi,c$)

CarriesSubjectBinding_s($chi,kappa,x,y_A,c$): the model route in chi contains or imports a live subject-binding claim presenting agent referent y_A under construal x as subject of the appraised target in $kappa$

SubjectTargetsOn_s(x,y_A,z_A,chi,c) $\leftrightarrow$ the subject-binding node for y_A -under- x has a directed path in ChainGraph(chi) to the formed target binding or appraisal node for z_A

TimeInOrAnchors($j,delta,t$) $\leftrightarrow$ TimeIn(j,t) or [not exists t' TimeIn(j,t') and TimeAnchorOf($delta$)= t]

A surface sentence may omit its subject while the derivation contains one. "What happened was wrong" can be agent-implicit when a route-carried subject binding links a person or institution to the event. By contrast, an appraisal such as "the earthquake was terrible" when the system represents no person, collective, institution, or other agent as subject of the event lacks MoralSubjectOf and cannot receive moral type. The moral relationship may still ground duties to rescue, mourn, repair, or prepare; it does not make the disaster an immoral agent.

4b.2 The moral-judgment biconditional

MoralJudgment($c,kappa$) $\leftrightarrow$ EvaluableFor_s($kappa,c$) and Outputs($c,kappa$)

and exists $K,chi,m,x,y_A,o,y,t,f,n,n_o,n_i$ [K =ComparisonClassOf($kappa,c$) and f =FrameOf(BaseOf($kappa$)) and
MoralRelationRep_s(m,x,y_A,o,y,t,f,K,c)

and MoralSubjectOf_s($kappa,x,y_A,chi,c$) and TimeInOrAnchors(Content($kappa$),BaseOf($kappa$), t)

and SubjectBindingNode_s(n,x,y_A,kappa,chi,c) and LocalActualContribution(n,RelevantEval(kappa),chi,kappa,c) and ObjectBindingNode_s(n_o,o,y,chi,BaseOf(kappa),c) and BoundProfileIntegrationOn_s(n_i,m,n,n_o,RelevantEval(kappa),chi,kappa,c)

and RelationshipProfileContributes_s(m,x,y_A,o,y,t,f,RelevantEval(kappa),chi,kappa,K,c) and BearsOnOn(kappa,o,y,RelevantEval(kappa),chi,c)]

The same actually contributing chain must carry and integrate the qualifying relationship, moral subject binding, and object binding at an independently typed relation-bearing evaluation input. The exact typed relationship-profile node, subject-binding node, and object-binding node must each locally contribute to the token-relative evaluation under the fixed constitutive scheme. Directed reachability is necessary but not sufficient. In particular, the node n that enters BoundProfileIntegrationOn_s is itself required to satisfy LocalActualContribution; a different subject-binding node cannot supply the contribution witness. A relationship in another frame, an unrepresented subject supplied by the analyst, a merely reachable but profile-silent node, or a moral object attached only after the output is known does not suffice.

FirstPersonal(c,kappa) $\leftrightarrow$ MoralJudgment(c,kappa) and exists chi,x,y_A[MoralSubjectOf_s(kappa,x,y_A,chi,c) and SelfReferent_s(y_A,c)]

ThirdPersonal(c,kappa) $\leftrightarrow$ MoralJudgment(c,kappa) and exists chi,x,y_A[MoralSubjectOf_s(kappa,x,y_A,chi,c) and not SelfReferent_s(y_A,c)]

First- and third-personality attach to the represented moral subject, not to grammatical surface form. A token with several moral subjects may receive a mixed classification; a token for which the system represents no agent as moral subject receives neither.

4b.3 One content, several provenances

MoralContent(c,j) $\leftrightarrow$ exists kappa[Outputs(c,kappa) and Content(kappa)=j and MoralJudgment(c,kappa)]

The same content may also have a nonmoral token. A principle, policy, or hygiene rule may produce the same words as a relationship-grounded judgment. Philosophical, legal, institutional, or moral vocabulary does not settle psychological type.

4b.4 Typed moral outputs

- A prospective token is moral when its represented subject stands in a qualifying relationship whose live provenance organizes GuideStruct.
- A violation token is moral when the same relational structure organizes the applicable requirement whose breach is represented.
- Wrongness and blame require an explicit or implicit represented subject; blame adds addressability, answerability, and requirement-specific responsibility.
- Praise is a favorable agent-act appraisal and therefore has an explicit represented subject.
- A generic Appraise token may target an event, outcome, object, relation, institution, or way of life. It is moral only when a represented moral subject is explicitly or implicitly linked to that target inside the same actually contributing chain.

Institutions and collectives are not exceptions. They receive moral type only when the judging system represents the institution itself, or relevant members acting through it, as possessing the practical and moral capacities required of a subject. A natural event with no represented agent can be tragic, harmful, sacred, beautiful, or terrible without being immoral.

4b.5 One provisional frame implementation

A frame-level score is one possible way to construct an outcome ordering. Action and outcome requirements, priority graphs, fallback structures, subject-target bindings, and appraisal structures require their own traceable construction. No constitutive thesis depends on one aggregation rule.

4c. Worked Composition Tests

The following traces are executable integration fixtures as well as explanatory examples. Each fixes an actual support graph and asks whether every stage produces the typed object consumed by the next. They demonstrate finite existence and interface compatibility; they do not establish that a real judgment has been empirically reconstructed correctly.

4c.1 Trace A: first-personal moral care

Context c contains a care frame. The judging system is the agent referent, presented under a parent construal; the child is the object referent, presented under a daughter construal. A spoilage representation supplies outcome-sensitive information. The represented parent-child relationship constructs a health requirement. One realized live chain in the token's actual support graph carries both the relationship and the child binding to the guide structure.

The relationship's positioning builder is intimacy-mediated in this fixture:

$\text{IntimacyRP}_s(\lambda_{\text{care}}, m_{\text{daughter}}, x_{\text{self}}, y_{\text{self}}, o_{\text{daughter}}, y_{\text{daughter}}, t, f_{\text{care}}, c)$, and therefore $\text{RelationalPositioning}_s$ holds for the same exact tuple. Intimacy is not required in other moral tokens.

$m_{\text{daughter}} \rightarrow u_{\text{health}} \rightarrow \text{Acc}_{\delta M} \rightarrow \text{Guide}_{\delta M} \rightarrow \kappa_M$

$\text{RealizedLiveChain}(\chi_M, \kappa_M, c) \wedge \text{OrganizesOn}(m_{\text{daughter}}, \text{GuideStruct}(\delta_M), \chi_M, \kappa_M, c)$

$\wedge \text{BearsOnOn}(\kappa_M, o_{\text{daughter}}, y_{\text{daughter}}, \text{GuideStruct}(\delta_M), \chi_M, c)$

The encoded ordering may independently favor discarding. The relationship-grounded acceptability boundary nevertheless performs evaluative work and is sufficient for moral type. First-personality follows from the self-referent condition; no separate attribution of moral vision to the judging system is required.

4c.2 Trace B: same content through a nonmoral route

Hold fixed the verbal content, feasible actions, spoilage representation, and action-outcome map, but use a generic hygiene or policy support graph. Its guide can still be $\{\text{discard}\}$. If no represented moral relationship reaches GuideStruct in any realized live chain, the token outputs but remains nonmoral.

$\text{Content}(\kappa_N) = \text{Content}(\kappa_M) = \text{OughtChoose}(x_{\text{self}}, \{\text{discard}\}, t)$

$\text{Outputs}(c, \kappa_N) \wedge \neg \text{MoralJudgment}(c, \kappa_N)$

This is the constructive witness for content-provenance nonidentity. The architecture does not infer moral cognition from a philosophical label, moral vocabulary, or behavior alone.

4c.3 Trace C: low attributed recognition - shark example

The inherent-valence, perceived-agency, and positioning profile, requirement, and support structure are held fixed. The represented nonself subject is attributed insufficient integrative-recognition capacity. The subject gate therefore fails. A shark is the example referent, but its entity kind is inert; the same low-recognition profile assigned to a human referent yields the same gate result. Answerability or requirement-specific responsibility may independently block blame in other cases, but they do not explain this fixture.

$\text{MV}_s(x_{\text{shark}}, o_{\text{child}}, t, f_{\text{attack}}, c) < \theta_M \rightarrow \neg \text{MoralSubjectGate}_s(x_{\text{shark}}, y_{\text{A_shark}}, o_{\text{child}}, t, f_{\text{attack}}, c)$

The occurrence can still receive a generic negative event or outcome appraisal. The model no longer forces condemnation of an occurrence into a blame token about an agent.

4c.4 Trace D: high attributed recognition - human example

The observer represents a competent stranger as standing toward the child through an operative complementary-role schema: capable responder / vulnerable patient. $\text{ComplementaryRoleIdentityRP}$ is registered for that exact stranger-child relationship; no standing-mediated residual category is used. This does not attribute intimacy, shared identity, empathy, endorsement, or positive valence to the stranger. The identity bridge and its downstream no-harm requirement are separately credited: the first realizes positioning; the second is a strict bender edge only if the exact core profile is preserved. Answerability, responsibility, and the nonself subject gate clear.

$m_{\text{stranger-child}} \rightarrow u_{\text{no-harm}} \rightarrow \text{Acc}_{\delta S} \rightarrow \text{Wrong}(x_{\text{stranger}}, a_{\text{attack}}, t)$

$\text{Wrong} \wedge \text{AnswerableTo} \wedge \text{ResponsibleCapacityFor} \rightarrow \text{Blameworthy}(x_{\text{stranger}}, a_{\text{attack}}, t)$

$\exists \chi_S [\text{RealizedLiveChain}(\chi_S, \kappa_S, c) \wedge \text{OrganizesOn}(m_{\text{stranger-child}}, \text{RelevantEval}(\kappa_S), \chi_S, \kappa_S, c)]$

$\wedge \text{BearsOnOn}(\kappa_S, o_{\text{child}}, y_{\text{child}}, \text{RelevantEval}(\kappa_S), \chi_S, c)]$

This trace states the substantive third-personal commitment openly: moral appraisal of another agent represents that agent as standing under a moral relation to the relevant object as construed by the judging system. The judged agent need not consciously recognize or endorse the relation.

4c.5 Trace E: violation-sensitive tragedy

Suppose every feasible act violates at least one non-overridden requirement. The model does not drop its deontic structure and revert automatically to outcome maximization. A represented fallback preorder first compares violation profiles and then uses the outcome ordering only as specified by that fallback.

$$\text{Acc_}\delta T = \emptyset \wedge \text{ViolationProfile_}\delta T(a_{\text{low}}) <_{\text{lex}} \text{ViolationProfile_}\delta T(a_{\text{high}})$$

$$\rightarrow \text{Guide_}\delta T = \{a_{\text{low}}\} \text{ even when } h_{\delta T}(a_{\text{high}}) >_{\text{out}} h_{\delta T}(a_{\text{low}})$$

$$a_{\text{low}} \in \text{Guide_}\delta T \wedge \text{WrongRelative}(a_{\text{low}}, \delta T) \wedge \text{OutputResidue}(a_{\text{low}}, u_{\text{low}}, \kappa_T, c)$$

4c.6 Trace F: agent-implicit appraisal and no represented agent

Surface content need not name the moral subject. In a negligence frame, token κ_N may have content $\text{Appraise}(\text{event_harm}, q_{\text{wrong}}, \text{negative}, t)$ - "what happened was wrong" - while its actually contributing chain carries a responsible official or institution as subject, the affected people as moral objects, and a subject-to-event link. MoralSubjectOf is satisfied implicitly, the relationship organizes the appraisal structure, and the token may be moral.

Now hold the negative event appraisal fixed but represent an earthquake without representing any agent as subject of the event. A generic Appraise token may still output from a tragic-harm dimension. Because no person, collective, institution, or other entity is represented as subject of the event, MoralSubjectOf fails. The appraisal is normatively or axiologically negative but not moral. The active relationship to victims may instead organize first-personal duties to rescue or rebuild.

$\text{Outputs}(c, \kappa_{\text{disaster}})$ and not exists x, y, A, chi $\text{MoralSubjectOf}_s(\kappa_{\text{disaster}}, x, y, A, \text{chi}, c) \rightarrow \text{not}$
 $\text{MoralJudgment}(c, \kappa_{\text{disaster}})$

4c.7 Trace G: moral self-relation and a downstream relation qualifier

A first-person token may bind the same thin self referent under role-distinct subject and object construals. The self disjunct clears only the subject gate; a moral self-relation requires the ordinary inherent-valence, agency-capacity, positioning, relationship-token, binding, access, and provenance conditions.

$\text{AgentConstrual}_s(x_{\text{self}}, y_{\text{self}}, f, c)$ and $\text{ObjectConstrual}_s(o_{\text{self}}, y_{\text{self}}, f, c)$
 $\text{MoralRelationRep}_s(m_{\text{self}}, x_{\text{self}}, y_{\text{self}}, o_{\text{self}}, y_{\text{self}}, t, f, \text{ComparisonClassOf}(\kappa_{\text{self}}, c), c)$ and
 $\text{RealizedLiveChain}(\text{chi}_{\text{self}}, \kappa_{\text{self}}, c)$

A represented ownership or self-sovereignty relation may then operate as a strict bender on a requirement, permission representation, priority ground, override warrant, residue condition, appraisal structure, or action description. A change to an agent or object construal, or to another core-individuating construal, is core-changing and must be typed separately rather than as a strict bender. Neither operation directly confers moral type. The finite companion witness represents permission-like structure through priority and override because its grammar does not yet contain a primitive permission judgment.

4c.8 Countermodel checks

Countermodel: Relationship and object paths occur in different chains

Expected result: Moral typing fails

Reason: MDT requires one shared realized chain.

Countermodel: A sole qualifying Relational Positioning witness is removed

Expected result: Moral typing fails

Reason: Existential typing has no remaining qualifying witness.

Countermodel: One qualifying positioning witness is removed while another remains

Expected result: Moral typing may survive

Reason: Deleting one witness does not defeat an existential claim supported by another witness.

Countermodel: Every qualifying positioning witness for the exact relationship is removed

Expected result: Moral typing fails

Reason: RelationalPositioningGate and the existential relationship witness both fail.

Countermodel: Relationship and exact bindings run on parallel same-chain subroutes that meet only at guide/output

Expected result: Moral typing fails

Reason: BoundProfileIntegrationOn requires convergence at a typed relation-bearing evaluation input.

Countermodel: A valence branch, agency, and relationship-strength branch pass but the exact relationship has no positioning

Expected result: Moral typing fails

Reason: RelationalPositioningGate is the non-scalar third-builder condition; access cannot substitute for it.

Countermodel: Valence criteria or configuration are registered for K1 while relationship-strength criteria or configuration are registered for K2, with $K1 \neq K2$

Expected result: Moral typing fails

Reason: One K must be bound to ComparisonClassOf(κ ,c) and threaded through every class-indexed relationship gate; registrations from distinct classes cannot be composed.

Countermodel: Either ordered positive/negative profile is replaced by a one-scalar summary

Expected result: Model instance invalid

Reason: Both the inherent-valence and derived relationship-strength profiles must survive through gates and aggregation.

Countermodel: A downstream duty is used to infer complementary-role identity

Expected result: Positioning not established

Reason: Identity bridge and downstream bender edge require separate evidence.

Countermodel: A derived requirement appears without construction ancestors or import certificate

Expected result: Route fails

Reason: The route is not provenance-closed.

Countermodel: One frame is duplicated under two labels

Expected result: No polyphony

Reason: Distinct event labels are insufficient without a derivationally relevant frame difference.

Countermodel: A rule, utility function, or policy is manipulated without a relationship path

Expected result: Normative output may survive; moral type fails

Reason: Philosophical or institutional labels do not confer provenance.

Countermodel: All acts are unacceptable and no fallback is represented

Expected result: No comparative guide

Reason: The theory does not silently erase constraints.

Countermodel: A merely competing consideration is treated as evidence

Expected result: No automatic block

Reason: Only admitted reliability, applicability, warrant, or target challenges enter the skeptical attack graph.

5. Modulators

Once a frame-indexed moral-relationship representation is active, other beliefs and representations can alter which evaluative structures and outputs it supports. Functional locus matters more than speed, consciousness, or emotional tone.

5.1 Builders, benders, attractors, and responsibility modifiers

Builders. Exactly three builder types determine whether the relationship profile is satisfied in a frame: (1) the inherent-valence profile with its governed branch boundary, (2) perceived agency with its floor, and (3) Relational Positioning with its existential gate. Agent and object construals are indices and binding conditions for those builders, not builder types. RelationshipStrengthProfile_s and its gate are derived valence-by-agency interaction machinery over the first two builders, not a fourth builder. The bound relationship token and profile-sensitive actual contribution are additional universal requirements, not further builders. Moral vision is separate and gates third-personal moral subjecthood.

Access conditions make routes functionally available but are not Relational Positioning. Identity-mediated positioning, intimacy, and direct encounter are the governed route families. Removing access blocks the affected contribution path. Removing one positioning witness removes that witness's organization while another witness may still satisfy RelationalPositioningGate. Removing the positioning builder means removing every qualifying witness for the exact relationship; the gate then fails.

Frame-relative constitution does not erase the builder-bender distinction. The three builders and token/contribution conditions determine whether the core relationship exists under *f*; the moral-subject gate determines whether a nonself agent can be morally appraised as subject to it. A *strict bender* presupposes and preserves that formed core while altering a downstream locus. A factor that changes a builder, a construal needed to individuate the core, or the bound relationship tuple is not a strict bender for that operation.

$\text{DownstreamBenderLocus}_s(z, \kappa, c) \leftrightarrow \text{ActionOutcomeLocus}_s(z, \kappa, c) \vee \text{ActionDescriptionLocus}_s(z, \kappa, c) \vee \text{RequirementLocus}_s(z, \kappa, c) \vee \text{PermissionLocus}_s(z, \kappa, c) \vee \text{ApplicabilityLocus}_s(z, \kappa, c) \vee \text{PriorityLocus}_s(z, \kappa, c) \vee \text{OverrideLocus}_s(z, \kappa, c) \vee \text{ResidueLocus}_s(z, \kappa, c) \vee \text{FeasibilityLocus}_s(z, \kappa, c) \vee \text{AlignmentLocus}_s(z, \kappa, c) \vee \text{AppraisalLocus}_s(z, \kappa, c) \vee \text{ResponsibilityAttributionLocus}_s(z, \kappa, c)$

$\text{PositiveInherentValueOn}_s(o, t, f, \text{chi}, \kappa, c)$: the positive inherent-value component actually carried on chain *chi* for token *kappa*.

$\text{NegativeInherentValueOn}_s(o, t, f, \text{chi}, \kappa, c)$: the negative inherent-value component actually carried on chain *chi* for token *kappa*.

$\text{AgencyCapOn}_s(x, o, t, f, \text{chi}, \kappa, c)$: the perceived-agency component actually carried on chain *chi* for token *kappa*.

$\text{RelationshipStrengthProfileOn}_s(x, o, t, f, \text{chi}, \kappa, c)$: the ordered positive and negative relationship-strength interaction profile actually carried on chain *chi* for token *kappa*.

$\text{CoreGateVerdictOn}_s(x, o, t, f, \text{chi}, \kappa, c)$: the component-gate verdict actually carried on chain *chi* for token *kappa*.

$\text{RelationalPositioningGateVerdictOn}_s(m, x, y_A, o, y, t, f, \text{chi}, \kappa, c)$: the Relational Positioning gate verdict actually carried by the exact structured path-witness organization on chain *chi* for token *kappa*.

$\text{RelationshipStrengthGateVerdictOn}_s(x, o, t, f, \text{chi}, \kappa, c)$: the derived relationship-strength gate verdict actually carried on chain *chi* for token *kappa*.

$\text{CoreMRVerdictOn}_s(m, x, y_A, o, y, t, f, \text{chi}, \kappa, c)$: the combined core and derived relationship-strength verdict actually carried on chain *chi* for token *kappa*.

$\text{BuilderQualifiedVerdictOn}_s(m, x, y_A, o, y, t, f, \text{chi}, \kappa, c)$: the verdict that the exact three-builder organization, including the chain-bound Relational Positioning witness, qualifies on chain *chi* for token *kappa*.

SubjectGateVerdictOn_s(x,y_A,t,f,chi,kappa,c): the moral-subject gate verdict actually carried on chain chi for token kappa.

SubjectGateBasisOn_s(B_subject,x,y_A,t,f,chi,kappa,c): B_subject is the exact registered basis label for self-reference or attributed moral recognition operative in the indexed subject-gate evaluation.

MoralVisionValueOn_s(V_subject,x,y_A,t,f,chi,kappa,c): V_subject is the exact represented moral-vision value operative in the indexed subject-gate evaluation.

MoralVisionFloorOn_s(F_subject,x,y_A,t,f,chi,kappa,c): F_subject is the exact registered moral-vision floor operative in the indexed subject-gate evaluation.

SubjectGateConfigurationOn_s(G_subject,x,y_A,t,f,chi,kappa,c): G_subject is the exact registered ordered record for the subject gate on chain chi for token kappa. None of its three values may be re-authored, omitted, or inferred merely from the Boolean subject-gate verdict when testing core-profile preservation.

SubjectGateConfigurationOn_s(G_subject,x,y_A,t,f,chi,kappa,c) ↔
 ∃B_subject,V_subject,F_subject[SubjectGateBasisOn_s(B_subject,x,y_A,t,f,chi,kappa,c) ∧
 MoralVisionValueOn_s(V_subject,x,y_A,t,f,chi,kappa,c) ∧ MoralVisionFloorOn_s(F_subject,x,y_A,t,f,chi,kappa,c) ∧
 G_subject=<B_subject,V_subject,F_subject>]

ValenceGateConfigurationOn_s(o,t,f,chi,kappa,c): the exact registered positive- and negative-component criteria operative for the valence gate on chain chi for token kappa under ComparisonClassOf(kappa,c).

AgencyGateConfigurationOn_s(x,o,t,f,chi,kappa,c): the exact registered agency floor, temporal-horizon domain, and response-affordance condition operative on chain chi for token kappa.

RelationshipStrengthGateConfigurationOn_s(x,o,t,f,chi,kappa,c): the exact registered positive- and negative-branch criteria operative for the derived relationship-strength gate on chain chi for token kappa under ComparisonClassOf(kappa,c).

BoundRelationalPositioningOn_s(m,x,y_A,o,y,t,f,chi,kappa,c): the exact bound Relational Positioning organization actually carried on chain chi for token kappa. Its identity is a structured record of chain, relationship, object- binding node, subject-binding node, governed route, explicit witness identifier when present, and, for identity-mediated positioning, every exact qualifying bridge schema, form, and subject/object binding. An opaque witness identifier alone, an identifier reused across routes, or a witness naming a nonexistent or different chain is not the same organization.

CoreProfileSnapshotOn_s(P_core,m,x,y_A,o,y,t,f,chi,kappa,c) ↔ RealizedLiveChain(chi,kappa,c) ∧ m∈Nodes(chi) ∧
 ∃G_subject[SubjectGateConfigurationOn_s(G_subject,x,y_A,t,f,chi,kappa,c) ∧ P_core=<chi,x,y_A,o,y,t,f,m,
 PositiveInherentValueOn_s(o,t,f,chi,kappa,c), NegativeInherentValueOn_s(o,t,f,chi,kappa,c),
 AgencyCapOn_s(x,o,t,f,chi,kappa,c), RelationshipStrengthProfileOn_s(x,o,t,f,chi,kappa,c),
 ValenceGateConfigurationOn_s(o,t,f,chi,kappa,c), AgencyGateConfigurationOn_s(x,o,t,f,chi,kappa,c),
 RelationshipStrengthGateConfigurationOn_s(x,o,t,f,chi,kappa,c), G_subject, CoreGateVerdictOn_s(x,o,t,f,chi,kappa,c),
 RelationshipStrengthGateVerdictOn_s(x,o,t,f,chi,kappa,c), RelationalPositioningGateVerdictOn_s(m,x,y_A,o,y,t,f,chi,kappa,c),
 CoreMRVerdictOn_s(m,x,y_A,o,y,t,f,chi,kappa,c), BuilderQualifiedVerdictOn_s(m,x,y_A,o,y,t,f,chi,kappa,c),
 SubjectGateVerdictOn_s(x,y_A,t,f,chi,kappa,c), BoundRelationalPositioningOn_s(m,x,y_A,o,y,t,f,chi,kappa,c)>]

NormalizedCoreProfilePayloadOn_s(Q_core,P_core,chi,kappa,c): Q_core is derived from the actual chain-indexed snapshot P_core by erasing only the chain label chi and by alpha-normalizing that same chain label inside its structured Relational Positioning witness records. Every other value remains exact: relationship, subject and object bindings; time and frame; positive and negative valence; agency value, horizon and affordances; all operative gate configurations; the derived relationship-strength profile; every component, core, relationship-strength, positioning, builder-qualified, and subject-gate verdict; and every route, node, explicit witness identifier, bridge schema, bridge form, and bridge binding. This is not an authored equality field.

BenderRemovalIntervention_s(I,u0,z,chi,kappa,c): a prospectively registered intervention that removes or varies only u0's tested operation on downstream locus z on chi, leaving all nonintervened support eligible for a chain counterpart comparison.

BenderRemovalCounterpartIdentity_s(I,u0,z,chi,chi_I,kappa,c): the registered removal identifier I names the same tested u0-to-z operation in the original chain registration and in counterpart chain chi_I; an unregistered, mismatched, or merely similar identifier does not satisfy this condition.

InterventionCounterpartTokenState_s(I,kappa[I,chi_I,c): a separately registered intervention-indexed counterpart token state supplies its own provenance-distinct relationship-witness and core-input record, event and case identity, time and frame, and named counterpart chain. Registered source and counterpart witness-input-state identifiers must bind distinct input records; direct reuse of the source witness object or its input-state identifier fails even when all semantic profile values are equal. Its core snapshot is derived from the counterpart inputs; an arbitrary token identifier or reuse of the source token's witness inputs does not satisfy this condition.

$\text{BenderRemovalMatch}_s(I, u_0, z, \text{chi}, \text{chi}_I, \text{kappa}, c) \leftrightarrow \text{BenderRemovalIntervention}_s(I, u_0, z, \text{chi}, \text{kappa}, c) \wedge$
 $\text{InterventionCounterpartChain}_s(I, \text{chi}, \text{chi}_I, \text{kappa}, c) \wedge \text{InterventionCounterpartTokenState}_s(I, \text{kappa}|I, \text{chi}_I, c) \wedge$
 $\text{BenderRemovalCounterpartIdentity}_s(I, u_0, z, \text{chi}, \text{chi}_I, \text{kappa}, c)$

$\text{CoreProfilePreservedOn}_s(u_0, m, x, y_A, o, y, t, f, z, \text{chi}, \text{kappa}, K, c) \leftrightarrow$
 $\exists I, \text{chi}_I, P_core_before, P_core_after, Q_core_before, Q_core_after [\text{BenderRemovalMatch}_s(I, u_0, z, \text{chi}, \text{chi}_I, \text{kappa}, c) \wedge$
 $K = \text{ComparisonClassOf}(\text{kappa}, c) \wedge K = \text{ComparisonClassOf}(\text{kappa}|I, c) \wedge$
 $\text{RelationshipProfileContributes}_s(m, x, y_A, o, y, t, f, \text{RelevantEval}(\text{kappa}), \text{chi}, \text{kappa}, K, c) \wedge$
 $\text{RelationshipProfileContributes}_s(m, x, y_A, o, y, t, f, \text{RelevantEval}(\text{kappa}|I), \text{chi}_I, \text{kappa}|I, K, c) \wedge$
 $\text{CoreProfileSnapshotOn}_s(P_core_before, m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c) \wedge$
 $\text{CoreProfileSnapshotOn}_s(P_core_after, m, x, y_A, o, y, t, f, \text{chi}_I, \text{kappa}|I, c) \wedge$
 $\text{NormalizedCoreProfilePayloadOn}_s(Q_core_before, P_core_before, \text{chi}, \text{kappa}, c) \wedge$
 $\text{NormalizedCoreProfilePayloadOn}_s(Q_core_after, P_core_after, \text{chi}_I, \text{kappa}|I, c) \wedge Q_core_before = Q_core_after]$

$\text{BenderPathIntervention}_s(J, u_0, z, E, \text{chi}, \text{kappa}, c)$: a prospectively registered intervention that varies the exact u_0 -to- z -to- E operation on chi rather than an unindexed effect of u_0 or z elsewhere in the graph.

$\text{BenderPathCounterpartIdentity}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c)$: the same registered path identifier J binds factor u_0 , typed locus z , relevant evaluation E , original chain chi , and counterpart chain chi_J . A registration for a parallel path, different locus, or different evaluation is not a match.

$\text{BenderPathEffectOn}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c)$: under matched path intervention J , the registered difference in E is attributable through the specified u_0 -to- z -to- E operation while declared nonpath drivers are preserved. The effect record must be validated against the actual named counterpart chain and intervention-indexed token; copying an authored effect Boolean or counterpart identifier without that chain does not satisfy this predicate.

$\text{BenderOperationActualOn}_s(u_0, z, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{RealizedLiveChain}(\text{chi}, \text{kappa}, c) \wedge u_0 \in \text{Nodes}(\text{chi}) \wedge z \in \text{Nodes}(\text{chi}) \wedge$
 $E \in \text{Nodes}(\text{chi}) \wedge E = \text{RelevantEval}(\text{kappa}) \wedge \text{Reach_chi}(u_0, z, c) \wedge \text{Reach_chi}(z, E, c) \wedge \text{AccessFor}_s(u_0, \text{chi}, \text{kappa}, c) \wedge$
 $\text{LocalActualContribution}(u_0, E, \text{chi}, \text{kappa}, c) \wedge \text{LocalActualContribution}(z, E, \text{chi}, \text{kappa}, c) \wedge$
 $\exists J, \text{chi}_J [\text{BenderPathIntervention}_s(J, u_0, z, E, \text{chi}, \text{kappa}, c) \wedge \text{InterventionCounterpartChain}_s(J, \text{chi}, \text{chi}_J, \text{kappa}, c) \wedge$
 $\text{InterventionCounterpartTokenState}_s(J, \text{kappa}|J, \text{chi}_J, c) \wedge \text{BenderPathCounterpartIdentity}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c) \wedge$
 $\text{BenderPathEffectOn}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c)]$

$\text{StrictBenderPathOn}_s(u_0, z, E, \text{chi}, \text{kappa}, c) \leftrightarrow \text{BenderOperationActualOn}_s(u_0, z, E, \text{chi}, \text{kappa}, c)$

$\text{BenderPathRelationshipContributionPreservedOn}_s(u_0, m, x, y_A, o, y, t, f, z, E, \text{chi}, \text{kappa}, K, c) \leftrightarrow$
 $K = \text{ComparisonClassOf}(\text{kappa}, c) \wedge \text{RelationshipProfileContributes}_s(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, K, c) \wedge$
 $\exists J, \text{chi}_J [\text{BenderPathIntervention}_s(J, u_0, z, E, \text{chi}, \text{kappa}, c) \wedge \text{InterventionCounterpartChain}_s(J, \text{chi}, \text{chi}_J, \text{kappa}, c) \wedge$
 $\text{InterventionCounterpartTokenState}_s(J, \text{kappa}|J, \text{chi}_J, c) \wedge \text{BenderPathCounterpartIdentity}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c) \wedge$
 $\text{BenderPathEffectOn}_s(J, u_0, z, E, \text{chi}, \text{chi}_J, \text{kappa}, c) \wedge K = \text{ComparisonClassOf}(\text{kappa}|J, c) \wedge$
 $\text{RelationshipProfileContributes}_s(m, x, y_A, o, y, t, f, \text{RelevantEval}(\text{kappa}|J), \text{chi}_J, \text{kappa}|J, K, c)]$

$\text{StrictBender}_s(u_0, m, x, y_A, o, y, t, f, z, \text{kappa}, K, c) \leftrightarrow K = \text{ComparisonClassOf}(\text{kappa}, c) \wedge \text{MoralRelationRep}_s(m, x, y_A, o, y, t, f, K, c) \wedge$
 $\text{DownstreamBenderLocus}_s(z, \text{kappa}, c) \wedge \exists E, \text{chi} [\text{StrictBenderPathOn}_s(u_0, z, E, \text{chi}, \text{kappa}, c) \wedge$
 $\text{BenderPathRelationshipContributionPreservedOn}_s(u_0, m, x, y_A, o, y, t, f, z, E, \text{chi}, \text{kappa}, K, c) \wedge$
 $\text{RelationshipProfileContributes}_s(m, x, y_A, o, y, t, f, E, \text{chi}, \text{kappa}, K, c) \wedge \text{BuilderQualifiedVerdictOn}_s(m, x, y_A, o, y, t, f, \text{chi}, \text{kappa}, c) \wedge$
 $\text{SubjectGateVerdictOn}_s(x, y_A, t, f, \text{chi}, \text{kappa}, c) \wedge \text{CoreProfilePreservedOn}_s(u_0, m, x, y_A, o, y, t, f, z, \text{chi}, \text{kappa}, K, c)]$

The path comparison and the separate removal comparison must each preserve actual contribution by the exact represented relationship profile in both the source and registered counterpart state. This requirement does not hold fixed Outputs, MoralJudgment, every binding-node contribution, or the final downstream result. Adding any of those stronger invariances would require a separate theoretical decision.

$\text{BuilderRealizer}_s(u_0, m, x, y_A, o, y, t, f, \text{kappa}, c) \vee \text{BuilderInput}_s(u_0, m, x, y_A, o, y, t, f, \text{kappa}, c) \vee$
 $\text{BuilderModulator}_s(u_0, m, x, y_A, o, y, t, f, \text{kappa}, c) \rightarrow \text{CoreChangingFactor}_s(u_0, m, x, y_A, o, y, t, f, \text{kappa}, c)$

BuilderRealizer_s identifies a representation or process that realizes a builder route or component. BuilderInput_s supplies information from which a builder component is constructed. $\text{BuilderModulator}_s$ changes the magnitude, form, or presence of a builder component. These types concern an operation, not an immutable property of the representation: the same representation may be a builder input on one edge and a strict bender on another.

The unindexed $\text{CoreChangingFactor}_s$ predicate says only that u_0 has at least one core-changing operation. It cannot exclude a separately evidenced strict-bender operation by that same support unit. The operation-specific exclusion remains enforced inside StrictBender_s by $\text{CoreProfilePreservedOn}_s$.

ValenceEdge_s(r,m,o,kappa,c), RelationalPositioningEdge_s(r,m,kappa,c), StrictBenderEdge_s(r,z,kappa,c), and AttractorEdge_s(r,f,c) are distinct functional credits. None implies another. One representation may bear several credits only when each edge independently satisfies its own binding, intervention, and contribution test.

IndependentEdgeCreditDiscipline_s(r,m,o,z,kappa,f,c) $\leftrightarrow$ PairwiseSeparatelyTested_s({ValenceEdge_s(r,m,o,kappa,c), RelationalPositioningEdge_s(r,m,kappa,c), StrictBenderEdge_s(r,z,kappa,c), AttractorEdge_s(r,f,c)},c)

Relation qualifiers. Some strict benders represent a further relation that qualifies the subject-object relation already presented in the frame. A relation qualifier may change what the represented relationship is taken to require, permit, prioritize, override, leave as residue, or appraise without becoming an additional universal builder.

RelationQualifierRep_s(r,k_{rel},x,y_A,o,y,t,f,c): active representation r presents relation-kind token k_{rel} as qualifying the represented relation between subject-under-x and object-under-o in frame f

QualifierOperatesOn_s(r,z,kappa,c) $\leftrightarrow \exists E, \chi$ StrictBenderPathOn_s(r,z,E, χ ,kappa,c)

QualifierOperatesOn_s establishes the registered path-specific operation through the typed downstream locus and relevant evaluation. It does not establish the same-chain relationship-profile contribution or exact-core preservation by itself. A relation qualifier receives strict-bender status only when the corresponding StrictBender_s clause, including the equal before/after core snapshots and matched removal/counterpart identity, also holds.

Admissible strict-bender loci include an action or outcome evaluation, action description, requirement, permission representation, applicability judgment, priority relation, override warrant, residue condition, alignment representation, appraisal target or structure, or responsibility attribution. A qualifier that instead changes an agent or object construal, inherent-valence component, AgencyCap, or Relational Positioning is typed as a core-changing factor for that edge. Which locus is active is a property of the represented route, not of the relation label alone.

RelationQualifierRep_s(r,...) does not by itself imply BuilderQualifiedMRCondition_s(m,x,y_A,o,y,t,f,K,c)

RelationQualifierRep_s(r,...) does not by itself imply MoralSubjectGate_s(x,y_A,o,t,f,c)

RelationQualifierRep_s(r,...) does not by itself imply MoralJudgment(c,kappa)

A qualifier can indirectly affect moral type when it has a separately credited core-changing edge. That operation is not a strict bender. No qualifier can bypass the moral-judgment biconditional.

Ownership as one relation qualifier. OwnershipConstrual_s(r,y_{owner},y_{owned},mu,t,f,c) implies a represented ownership qualifier. The mode mu is frame-relative and may be mixed. Nonexhaustive possibilities include stewardship, which may contribute special requirements or care priorities; sovereignty, which may contribute a represented permission, priority ground, licensed override, or residue condition; and absolute dominion, which may alter requirement applicability, suppress represented independent standing, shift the appraisal target toward the owner, or alter the object construal itself. The last operation is core-changing, not strict bending.

Ownership has no fixed sign. It may add represented responsibility, support a claimed license, or otherwise alter the represented relation. Legal title, possessive grammar, physical control, intimacy, and represented ownership are not interchangeable. A self-directed moral relation may coexist with a self-sovereignty qualifier. NDT thereby describes the represented permission or override without entailing that it is correct, universally available, or sufficient to defeat every self-directed requirement.

An ownership ideology organized around dominion can likewise be represented as combining selective duties with claimed control, reduced represented independent standing, or target shift. Antebellum slaveholding frames provide one historical example. Formal representation describes the frame's cognitive organization and does not determine whether its associated claims are correct, justified, or true.

Candidate bender sources from Moral Space (2021). Vulnerability, power, and sentience are retained as candidate sources of downstream bending. Their mention does not establish that they are nonmoral, psychometrically distinct, causally active, or strict benders in a given token. Each requires a separately identified downstream edge with the core profile held fixed. When one changes inherent valence, perceived agency, Relational Positioning, or a core construal, that operation is typed as builder realization, input, or modulation instead.

Attractors. An attractor changes the prior salience or accessibility of frames without entering the selected frame as a ranking-, standard-, priority-, alignment-, or appraisal-relevant premise.

Attractor(r,c) $\rightarrow$ BiasesSalience(r,Frames(c))

Responsibility modifiers. Representations of control, knowledge, coercion, opportunity, and agent-relative positioning may change blame or praise attributed to an agent construal without altering the underlying object visibility or overall act acceptability.

5.2 Contribution, override, and undercutting

Three routes must remain distinct. A new consequence, weight, or standard contributes to the evaluative structure. A priority ground contributes to an override while leaving the lower requirement applicable and violated. An undercutter challenges a premise, an inferential link, or a whole target. A high-stakes reason for lying does not make the truth-telling requirement inapplicable; it may license override. Evidence that a testimony is fabricated may undercut the premise, while evidence that a valid premise does not support the conclusion in this context may undercut the link.

5.3 Illustration: spoiled milk

A biohacker's claim that a small dose of spoiled milk improves health changes the represented consequence map and perhaps the ordering; it bends. A ritual identity can attract a sacred frame, and doctrine inside that frame can also bend. Evidence that the biohacker fabricated the data undercuts the credibility premise. Evidence that even accurate microbiome data do not warrant this dosage under the child's medical conditions undercuts an inferential link. A represented emergency that makes a normally applicable prohibition worth overriding contributes to the priority graph rather than undercutting the prohibition. The worked audit in Section 4c adds a separate point: the spoilage representation can support the same discard prescription through a nonmoral safety route, while a moral care token requires the daughter relationship to ground an ordering, requirement, priority, acceptability boundary, or appraisal component in the actual chain.

5.4 Robustness ordering

A working empirical hypothesis is Builders > Benders > Attractors. Builders may be comparatively perception-like because they constitute the relationship itself. Benders may be more stable than attractors because they alter content rather than salience. Responsibility modifiers, priority graphs, alignment models, and undercutters need not occupy one fixed position. None of these robustness claims is constitutive.

5.5 Institutional norms and moral zombies

Roles, rules, promises, and conventions may bend, attract, contribute requirements, or supply priority structures. They may also generate normative outputs with no moral grounding. Rule compliance motivated only by punishment, habit, authority, or script completion remains nonmoral when no represented moral relationship organizes the token-relative evaluative structure along a connected live chain. The historical moral origin of an institution does not moralize every later compliance token.

6. The Polyphony Principle

Polyphony concerns several indexed outputs, not a contradiction inside one undifferentiated premise set. The alignment layer is canonical rather than merely illustrative: without a surviving bridge, outputs are plural or incomparable, not automatically conflicting.

6.1 Mode-indexed admissibility and identity

$\text{ReconstructivelyAdmissible}_s(f, \kappa, c) \leftrightarrow \text{exists } \delta, x, t [\text{BaseOf}(\kappa) = \delta \text{ and } \text{FrameOf}(\delta) = f \text{ and } \text{ReconstructionFormed}_s(c, \delta, x, t)]$

$\text{FrameAdmissible}_s(f, \kappa, c) \leftrightarrow [\text{ModeOf}(\kappa) = \text{prospective and } \text{CurrentFrameFormed}(f, c)] \text{ or } [\text{ModeOf}(\kappa) = \text{reconstructive and } \text{ReconstructivelyAdmissible}_s(f, \kappa, c)]$

$\text{SameCase}_s(\kappa_1, \kappa_2, c) \leftrightarrow \text{exists } \zeta [\text{CaseOf}(\kappa_1) = \zeta \text{ and } \text{CaseOf}(\kappa_2) = \zeta \text{ and } \text{ConstruesEpisode}_s(\text{FrameOf}(\text{BaseOf}(\kappa_1)), \text{EpisodeConstrualOf}(\kappa_1), \zeta, c) \text{ and } \text{ConstruesEpisode}_s(\text{FrameOf}(\text{BaseOf}(\kappa_2)), \text{EpisodeConstrualOf}(\kappa_2), \zeta, c)]$

$\text{PluralPair}_s(\kappa_1, \kappa_2, c) \leftrightarrow \text{Outputs}(c, \kappa_1) \text{ and } \text{Outputs}(c, \kappa_2) \text{ and not } \text{SameJudgmentToken}(\kappa_1, \kappa_2, c) \text{ and not } \text{SameFrameToken}(\text{FrameOf}(\text{BaseOf}(\kappa_1)), \text{FrameOf}(\text{BaseOf}(\kappa_2)), c) \text{ and } \text{SameCase}_s(\kappa_1, \kappa_2, c) \text{ and } \text{FrameAdmissible}_s(\text{FrameOf}(\text{BaseOf}(\kappa_1)), \kappa_1, c) \text{ and } \text{FrameAdmissible}_s(\text{FrameOf}(\text{BaseOf}(\kappa_2)), \kappa_2, c) \text{ and } \text{FrameDifferenceRelevant}(\text{FrameOf}(\text{BaseOf}(\kappa_1)), \text{FrameOf}(\text{BaseOf}(\kappa_2)), \kappa_1, \kappa_2, c)$

6.2 Action, target, and dimension alignment

$MentionedActs(\kappa) :=$ the frame-relative action descriptions occurring in $Content(\kappa), Guide(BaseOf(\kappa))$, or the appraised scenario

$AlignmentRoute_s(\rho, \beta, \delta_1, \delta_2, y_A, t, c) \leftrightarrow Route(\rho, \beta, c)$ and carries both agent-binding claims and the shared temporal anchor and constructs maps from each action space into common execution space $X_{\{y_A, t, c\}}$

$SurvivesOn(\rho, \beta, c) \leftrightarrow LiveRoute(\rho, \beta, c)$

$Aligned(\beta, \delta_1, \delta_2, c) \leftrightarrow$ exists $\rho, y_A, t [AlignmentRoute_s(\rho, \beta, \delta_1, \delta_2, y_A, t, c)$ and $SurvivesOn(\rho, \beta, c)]$

$AlignmentCovers(\beta, \kappa_1, \kappa_2, c) \leftrightarrow$ [for every a in $MentionedActs(\kappa_1)$, exists α $Realizes_beta(\alpha, a, BaseOf(\kappa_1), c)$] and [for every a in $MentionedActs(\kappa_2)$, exists α $Realizes_beta(\alpha, a, BaseOf(\kappa_2), c)$]

$PairAligned(\beta, \kappa_1, \kappa_2, c) \leftrightarrow Aligned(\beta, BaseOf(\kappa_1), BaseOf(\kappa_2), c)$ and $AlignmentCovers(\beta, \kappa_1, \kappa_2, c)$

If $AlignmentCovers$ fails, the pair is unaligned. Empty execution extensions created by missing maps never witness conflict.

$TargetAlignmentModel(\tau, z_1, \delta_1, z_2, \delta_2, c) \leftrightarrow \tau$ maps the formed target construals z_1, z_2 into one represented comparison target and has a surviving route

$TargetAligned(\tau, \kappa_1, \kappa_2, c) \leftrightarrow TargetAlignmentModel(\tau, AppraisalTargetOf(\kappa_1), BaseOf(\kappa_1), AppraisalTargetOf(\kappa_2), BaseOf(\kappa_2), c)$ and $Survives(\tau, c)$

$DimensionAlignmentModel(\gamma, q_1, \delta_1, q_2, \delta_2, c) \leftrightarrow \gamma$ maps X_{δ_1, q_1} and X_{δ_2, q_2} into a common appraisal space with pre-specified order-preservation conditions and a surviving route

$DimAligned(\gamma, q_1, \delta_1, q_2, \delta_2, c) \leftrightarrow DimensionAlignmentModel(\gamma, q_1, \delta_1, q_2, \delta_2, c)$ and $Survives(\gamma, c)$

6.3 Conflict for indexed tokens

$Exec_beta(C, \delta, c) := \{\alpha \text{ in } X_{\{AgentReferentOf(\delta), TimeAnchorOf(\delta), c\}} : \text{exists } a \text{ in } C \text{ } Realizes_beta(\alpha, a, \delta, c)\}$

$ChoiceSetOf(\kappa) = C \leftrightarrow$ exists $x, t [Content(\kappa) = OughtChoose(x, C, t)]$

$Condemnatory(\kappa, x, a, t) \leftrightarrow Content(\kappa) = Wrong(x, a, t)$ or $Content(\kappa) = Blameworthy(x, a, t)$ or exists u $Content(\kappa) = Violation(x, a, u, t)$

$CondemnedExec_beta(\kappa, \delta, c) := \{\alpha : \text{exists } x, a, t [Condemnatory(\kappa, x, a, t) \text{ and } Realizes_beta(\alpha, a, \delta, c)]\}$

$AcceptableExec_beta(\delta, c) := Exec_beta(Acc_delta, \delta, c)$

$ProspectiveConflict(\kappa_1, \kappa_2, c) \leftrightarrow$ exists $\beta, C_1, C_2 [ChoiceSetOf(\kappa_1) = C_1$ and $ChoiceSetOf(\kappa_2) = C_2$ and $PairAligned(\beta, \kappa_1, \kappa_2, c)$ and $Exec_beta(C_1, BaseOf(\kappa_1), c)$ not empty and $Exec_beta(C_2, BaseOf(\kappa_2), c)$ not empty and $Exec_beta(C_1, BaseOf(\kappa_1), c) \cap Exec_beta(C_2, BaseOf(\kappa_2), c) = \emptyset]$

$GuideCondemnationConflict(\kappa_g, \kappa_w, c) \leftrightarrow$ exists $\beta, C_g [ChoiceSetOf(\kappa_g) = C_g$ and $PairAligned(\beta, \kappa_g, \kappa_w, c)$ and $Exec_beta(C_g, BaseOf(\kappa_g), c) \cap CondemnedExec_beta(\kappa_w, BaseOf(\kappa_w), c)$ not empty and $Exec_beta(C_g, BaseOf(\kappa_g), c) \cap AcceptableExec_beta(BaseOf(\kappa_w), c) = \emptyset]$

$ActIn(j, a)$: content j appraises frame-relative action description a

$Appraises_beta(\kappa, \alpha, c) \leftrightarrow$ exists $a [ActIn(Content(\kappa), a)$ and $Realizes_beta(\alpha, a, BaseOf(\kappa), c)]$

$SameAct_beta(\kappa_1, \kappa_2, c) \leftrightarrow$ exists $\alpha [Appraises_beta(\kappa_1, \alpha, c)$ and $Appraises_beta(\kappa_2, \alpha, c)]$

$PositiveAppraisal(\kappa) \leftrightarrow$ exists $x, a, t [Content(\kappa) = Praiseworthy(x, a, t)]$ or exists $z_A, q, t [Content(\kappa) = Appraise(z_A, q, positive, t)]$

$NegativeAppraisal(\kappa) \leftrightarrow$ exists $x, a, t [Content(\kappa) = Wrong(x, a, t)$ or $Content(\kappa) = Blameworthy(x, a, t)]$ or exists $z_A, q, t [Content(\kappa) = Appraise(z_A, q, negative, t)]$

$PolarAppraisal(\kappa) \leftrightarrow PositiveAppraisal(\kappa)$ or $NegativeAppraisal(\kappa)$

$PolarityOf(\kappa) = +1 \leftrightarrow PositiveAppraisal(\kappa)$

$PolarityOf(\kappa) = -1 \leftrightarrow NegativeAppraisal(\kappa)$

ActionLikeAppraisal(κ) $\leftrightarrow$ exists x, a, t [Content(κ)= Praiseworthy(x, a, t) or Content(κ)=Wrong(x, a, t) or Content(κ)=Blameworthy(x, a, t)]

FormedNonActionAppraisal(κ, c) $\leftrightarrow$ exists z_A, q, pol, t [Content(κ)=Appraise(z_A, q, pol, t) and AssessmentTargetFormed_ $s(c, \kappa, BaseOf(\kappa), z_A, t)$ and TargetSort(z_A) not in {action, agent-act}]

ActionAppraisalPair(κ_1, κ_2) $\leftrightarrow$ ActionLikeAppraisal(κ_1) and ActionLikeAppraisal(κ_2)

NonActionAppraisalPair(κ_1, κ_2, c) $\leftrightarrow$ FormedNonActionAppraisal(κ_1, c) and FormedNonActionAppraisal(κ_2, c)

SameAppraisalTargetOn($\tau, \kappa_1, \kappa_2, c$) $\leftrightarrow$ τ aligns AppraisalTargetOf(κ_1) and AppraisalTargetOf(κ_2) to one represented comparison target

SameOverallQuestion(κ_1, κ_2) $\leftrightarrow$ exists Q [QuestionOf(κ_1)= Q and QuestionOf(κ_2)= Q]

ComparableAppraisalTarget(κ_1, κ_2, c) $\leftrightarrow$ [ActionAppraisalPair(κ_1, κ_2) and exists β [PairAligned($\beta, \kappa_1, \kappa_2, c$) and SameAct_ $\beta(\kappa_1, \kappa_2, c)$]] or [NonActionAppraisalPair(κ_1, κ_2, c) and exists τ [TargetAligned($\tau, \kappa_1, \kappa_2, c$) and SameAppraisalTargetOn($\tau, \kappa_1, \kappa_2, c$)]]

OverallAppraisalConflict(κ_1, κ_2, c) $\leftrightarrow$ PolarAppraisal(κ_1) and PolarAppraisal(κ_2) and SameOverallQuestion(κ_1, κ_2) and PolarityOf(κ_1)= $-PolarityOf(\kappa_2)$ and ComparableAppraisalTarget(κ_1, κ_2, c)

AspectAppraisalConflict(κ_1, κ_2, c) $\leftrightarrow$ PolarAppraisal(κ_1) and PolarAppraisal(κ_2) and exists q_1, q_2, γ [LocalDimensionOf(κ_1)= q_1 and LocalDimensionOf(κ_2)= q_2 and DimAligned($\gamma, q_1, BaseOf(\kappa_1), q_2, BaseOf(\kappa_2), c$) and OpposedUnder($\gamma, \kappa_1, \kappa_2, c$) and ComparableAppraisalTarget(κ_1, κ_2, c)]

GenericTargetConflict(κ_1, κ_2, c) $\leftrightarrow$ exists $z_1, q_1, p_1, t_1, z_2, q_2, p_2, t_2, \tau$ [Content(κ_1)=Appraise(z_1, q_1, p_1, t_1) and Content(κ_2)=Appraise(z_2, q_2, p_2, t_2) and PolarityOf(κ_1)= $-PolarityOf(\kappa_2)$ and TargetAligned($\tau, \kappa_1, \kappa_2, c$) and SameAppraisalTargetOn($\tau, \kappa_1, \kappa_2, c$) and [SameOverallQuestion(κ_1, κ_2) or exists γ DimAligned($\gamma, q_1, BaseOf(\kappa_1), q_2, BaseOf(\kappa_2), c$)]]

Conflict(κ_1, κ_2, c) $\leftrightarrow$ ProspectiveConflict(κ_1, κ_2, c) or GuideCondemnationConflict(κ_1, κ_2, c) or GuideCondemnationConflict(κ_2, κ_1, c) or OverallAppraisalConflict(κ_1, κ_2, c) or AspectAppraisalConflict(κ_1, κ_2, c) or GenericTargetConflict(κ_1, κ_2, c)

Violation remains nonpolar at the overall level. A natural-disaster appraisal may conflict with another appraisal in a general evaluative space without thereby becoming moral; moral type and conflict type are separate questions.

6.4 Plurality, incomparability, and polyphony

OughtToken(κ) $\leftrightarrow$ exists x, C, t [Content(κ)=OughtChoose(x, C, t)]

CondemnatoryToken(κ) $\leftrightarrow$ exists x, a, t Condemnatory(κ, x, a, t)

PracticalPair(κ_1, κ_2) $\leftrightarrow$ [OughtToken(κ_1) and OughtToken(κ_2)] or [OughtToken(κ_1) and CondemnatoryToken(κ_2)] or [CondemnatoryToken(κ_1) and OughtToken(κ_2)]

ActionComparisonReady(κ_1, κ_2, c) $\leftrightarrow$ PracticalPair(κ_1, κ_2) and exists β [PairAligned($\beta, \kappa_1, \kappa_2, c$)]

OverallAppraisalReady(κ_1, κ_2, c) $\leftrightarrow$ PolarAppraisal(κ_1) and PolarAppraisal(κ_2) and SameOverallQuestion(κ_1, κ_2) and ComparableAppraisalTarget(κ_1, κ_2, c)

AspectAppraisalReady(κ_1, κ_2, c) $\leftrightarrow$ exists q_1, q_2, γ [LocalDimensionOf(κ_1)= q_1 and LocalDimensionOf(κ_2)= q_2 and DimAligned($\gamma, q_1, BaseOf(\kappa_1), q_2, BaseOf(\kappa_2), c$) and ComparableAppraisalTarget(κ_1, κ_2, c)]

GenericTargetReady(κ_1, κ_2, c) $\leftrightarrow$ exists $z_1, q_1, p_1, t_1, z_2, q_2, p_2, t_2, \tau$ [Content(κ_1)=Appraise(z_1, q_1, p_1, t_1) and Content(κ_2)=Appraise(z_2, q_2, p_2, t_2) and TargetAligned($\tau, \kappa_1, \kappa_2, c$) and SameAppraisalTargetOn($\tau, \kappa_1, \kappa_2, c$) and [SameOverallQuestion(κ_1, κ_2) or exists γ DimAligned($\gamma, q_1, BaseOf(\kappa_1), q_2, BaseOf(\kappa_2), c$)]]

ComparisonReady(κ_1, κ_2, c) $\leftrightarrow$ ActionComparisonReady(κ_1, κ_2, c) or OverallAppraisalReady(κ_1, κ_2, c) or AspectAppraisalReady(κ_1, κ_2, c) or GenericTargetReady(κ_1, κ_2, c)

UnalignedPair_ $s(\kappa_1, \kappa_2, c) \leftrightarrow$ PluralPair_ $s(\kappa_1, \kappa_2, c)$ and not ComparisonReady(κ_1, κ_2, c)

Unaligned plurality has more than one diagnostic form. A pair may lack any represented action or target bridge, or it may share an execution/target bridge while still lacking the shared overall question or surviving dimension map required for appraisal comparison. These are explanatory subtypes of `UnalignedPair`, not additional top-level formal outcomes. Shared execution therefore does not by itself imply that two appraisals are comparison-ready.

$\text{PolyphonicPair}_s(\kappa_1, \kappa_2, c) \leftrightarrow \text{PluralPair}_s(\kappa_1, \kappa_2, c)$ and $\text{Coherent}_s(\text{FrameOf}(\text{BaseOf}(\kappa_1)), \kappa_1, c)$ and $\text{Coherent}_s(\text{FrameOf}(\text{BaseOf}(\kappa_2)), \kappa_2, c)$ and $\text{Conflict}(\kappa_1, \kappa_2, c)$

$\text{Polyphony}(c) \leftrightarrow \text{exists } \kappa_1, \kappa_2 \text{ PolyphonicPair}_s(\kappa_1, \kappa_2, c)$

Prospective, reconstructive, and mixed-mode polyphony are distinguished by `ModeOf`. Higher-order arbitration is optional: an arbitration token may itself enter further plurality or conflict, and no global closure is assumed.

6.5 Illustrations and limits

Raskolnikov-style appraisals may differ in agent construal, episode construal, action description, motive, object scope, local dimension, and overall question while remaining anchored to one broad case. The theory classifies them as polyphonic only where the relevant bridges survive and the outputs conflict. Where no bridge is represented, the correct result is unaligned plurality.

7. Connected Route Survival and Grounded Defeat

NDT distinguishes contributions, overrides, and challenges. A consideration that changes an ordering, action map, requirement, priority, feasibility representation, or appraisal contributes. A reason for violating an applicable requirement belongs in the priority structure. An undercutter challenges the reliability, applicability, or warrant of a premise, link, import, or target.

7.1 Contributions and overrides are not undercutters

`Contributes(d,z,f,c)`: d helps construct or alter z inside f

`OverrideGround(v,u,a, $\delta$ ,c)`: v is represented as a reason to violate applicable u in choosing a

`ChallengeEligible(d,u,c)`: d is independently identifiable, active, targeted to u, and represented as bearing on reliability, truth, applicability, source, availability, warrant, import, or target formation

Contrary evidence is not automatically an undercutter. It may contribute to a different estimate, ordering, frame, or token. Only a consideration represented as challenging a premise or inferential connection enters the attack layer.

7.2 Structural routes, exports, and parent-child imports

`Basis(z)` records traceable premise-like nodes represented as possible support for target z. `Warrants(z)` records represented mappings, inference rules, aggregation operations, and transition links. Routes are derivational graphs, not bags of conclusions.

$\text{Route}(\rho, z, c) \leftrightarrow \rho \text{ subset of } \text{SupportGraphOf}(z, c) \text{ and } \text{DerivationallySufficient}(\rho, z, c) \text{ and } \text{ProvClosed}(\rho, c) \text{ and } \text{MinimalDerivation}(\rho, z, c)$

$\text{ProvClosed}(\rho, c) \leftrightarrow \text{for every derived } u \text{ in } \text{Units}(\rho), \text{ each immediate construction ancestor is either in } \rho \text{ or is supplied by an explicit ImportsAs edge from the declared parent route}$

$\text{NativeYield}(\rho, c) := \{w : \text{ConstructsOrLicenses}(\rho, w, c)\}$

$\text{ReExport}(\rho, c) := \{w : w \text{ was validly imported into } \rho \text{ and } \rho \text{ explicitly exposes } w \text{ to a child}\}$

$\text{Export}(\rho, c) := \text{NativeYield}(\rho, c) \text{ union } \text{ReExport}(\rho, c)$

$\text{ImportsAs}(w, u, \rho_p, \rho_c, c) \rightarrow \text{ParentRoute}(\rho_c) = \rho_p \text{ and } w \text{ in } \text{Export}(\rho_p, c) \text{ and } u \text{ in } \text{Units}(\rho_c)$

$\text{Anchors}(\rho_c, \rho_p, c) \leftrightarrow \text{ParentRoute}(\rho_c) = \rho_p \text{ and every inherited unit in } \rho_c \text{ is supplied by a valid immediate import from } \rho_p \text{ and at least one such import occurs}$

No free-standing provenance certificate is assumed. Compressed implementations may use certificates internally, but the formal theory requires their content to reduce to an identifiable source route, export product, child unit, ancestry relation, and defeat-sensitive import edge.

7.3 Attack graph and grounded labeling

Only independently admitted challenges enter the skeptical attack graph. Admission is representational and may be automatic, affective, implicit, motivated, or deliberative; it need not be voluntary.

$\text{AdmittedChallenge}_s(d,u,c) \leftrightarrow \text{Registered}_s(d,c) \text{ and } \text{Targets}_s(d,u,c) \text{ and } \text{Credibility}_s(d,c) \geq \text{theta_cred} \text{ and } \text{ChallengeRelevance}_s(d,u,c) \geq \text{tau_chal}$

$\text{PremiseUndercuts}(d,p,\rho,c) \rightarrow \text{AdmittedChallenge}_s(d,p,c)$

$\text{LinkUndercuts}(d,\ell,\rho,c) \rightarrow \text{AdmittedChallenge}_s(d,\ell,c)$

$\text{DirectUndercuts}(d,z,c) \rightarrow \text{AdmittedChallenge}_s(d,z,c)$

$\text{Attacks}(d,d',c) \leftrightarrow \text{MetaUndercuts}(d,d',c)$

Grounded labeling supplies the reference skeptical implementation. An admitted challenge is IN when all its admitted attackers are OUT; it is OUT when attacked by an IN challenge; otherwise it remains UNDECIDED. Undecided challenges are not treated as defeated. Credulous, preferred, graded, or probabilistic semantics remain compatible implementation families.

Output change is not part of the definition of admission. Registration, target recognition, credibility, and challenge relevance must be measured independently if the defeat layer is to be falsifiable.

7.4 Premise, link, and target availability

$\text{PremiseAvailable}(p,\rho,c) \leftrightarrow \text{ActiveIn}(p,\rho,c) \wedge \forall d[\text{PremiseUndercuts}(d,p,\rho,c) \rightarrow \text{Defeated}(d,c)]$

$\text{LinkAvailable}(\ell,\rho,c) \leftrightarrow \text{LicensedIn}(\ell,\rho,c) \wedge \forall d[\text{LinkUndercuts}(d,\ell,\rho,c) \rightarrow \text{Defeated}(d,c)]$

$\text{DirectClear}(z,c) \leftrightarrow \forall d[\text{DirectUndercuts}(d,z,c) \rightarrow \text{Defeated}(d,c)]$

The admission rule is load-bearing. A merely noticed contrary consideration does not silence a judgment unless the system represents it as a sufficiently strong challenge to a specified premise, warrant, import, or target.

7.5 Live routes, connected chains, survival, and output

$\text{FunctionallyActiveRoute}(\rho,z,c) \leftrightarrow \text{Route}(\rho,z,c) \text{ and the route's processes occur or remain functionally active in the modeled episode}$

$\text{LiveRoute}(\rho,z,c) \leftrightarrow \text{FunctionallyActiveRoute}(\rho,z,c) \text{ and } \text{DirectClear}(z,c) \text{ and every premise and link in } \rho \text{ is available}$

$\text{FunctionallyActiveChain}(\chi_i,\kappa,c) \leftrightarrow \text{the frame, model, and token routes are functionally active and correctly anchored}$

$\text{LiveChain}(\chi_i,\kappa,c) \leftrightarrow \text{the frame, model, and token routes are live and the model route is anchored to the frame route and the token route to the model route}$

$\text{RealizedLiveChain}(\chi_i,\kappa,c) \leftrightarrow \text{LiveChain}(\chi_i,\kappa,c) \text{ and } \text{ActualContribution}(\chi_i,\kappa,c)$

$\text{Outputs}(c,\kappa) \leftrightarrow \text{EvaluableFor}_s(\kappa,c) \text{ and } \text{Candidate}(c,\kappa) \text{ and exists } \chi_i \text{ RealizedLiveChain}(\chi_i,\kappa,c)$

A route is not actual merely because it can be reconstructed from $\text{Basis}(\kappa)$. Under redundant support, several chains may count as actual when each is active and belongs to an actual minimal sufficient set for the pre-specified token profile. If empirical evidence does not distinguish active support from available backup, provenance may remain unresolved.

Activity, survival, and contribution are therefore distinct. An active route may be challenged or defeated and fail to be live; a live route may still fail the fixed-scheme contribution test; only a surviving, actually contributing chain can support output.

7.6 Why connected survival adds genuine work

Connected survival prevents support laundering across levels and across possible routes. A model may survive through one action map while a token imports another defeated map; a relationship may organize evaluation on one chain while the object binding occurs only on another; a requirement may remain available while the warrant from it to blame fails. Each failure is represented without pretending that a premise is false.

7.7 Why the angel usually bends rather than defeats

A new consequence normally changes the action-outcome map, ordering, requirement, priority, or fallback and therefore contributes. Evidence that the milk is fresh undercuts a premise; evidence that an angelic warning is hallucinated undercuts

that premise; evidence that even a truthful warning does not support the probability estimate undercuts a warrant. The categories are functional, not rhetorical.

7.8 Anti-vacuity constraint

A challenge attribution is permissible only when the consideration, target, admission threshold, and attack relation are identifiable independently of the desired output. A support graph, relationship path, implicit subject, frame boundary, control aggregation, or fallback ordering may not be invented after the fact to explain a classification.

The same anti-post-hoc rule applies to non-admission. A failed output-suppression prediction cannot be rescued merely by asserting that the challenge was not admitted; registration, target recognition, credibility, and relevance must be measured without using the output as their criterion.

7.9 Why grounded skepticism is retained

Grounded semantics gives a unique skeptical extension and is useful for conservative model checking. It is retained as the reference implementation, not declared the universal psychology of evidential conflict. Empirical work must identify which considerations a system admits as challenges and whether its practical cognition is skeptical, credulous, graded, probabilistic, or frame-splitting. The executable companion implements the skeptical baseline and can be extended with alternatives.

Survival is an internal modeled status under the selected challenge semantics. It is not objective warrant, truth, or justification. Under the skeptical reference, an admitted undefeated or undecided challenge blocks the targeted support; other resolvers remain replaceable implementations rather than v52 commitments.

8. Evidence, adequacy, and empirical development

Four evaluative levels must be kept apart. Comparative explication adequacy asks whether the proposed kind is noncircular, preserves paradigm and contrast cases, integrates with general normativity, and earns discriminating consequences. Formal countermodel structure asks whether the architecture states finite conditions under which a registered formal claim fails. Operational identification asks whether latent variables and rival organizations can be distinguished independently of the target classification. Calibrated measurement requires validated indicators, weights, thresholds, uncertainty, and a declared population or operating range. Success at one level does not establish success at another. Practical tractability - the feasibility, fidelity, and burden of available interventions and measurements - cuts across all four; present intractability is a limitation, not confirmation.

Falsifiability is neither the sole nor the central measure of NDT's value. The theory's primary contribution is to specify a neglected explanandum, distinguish moral from nonmoral normativity without circular reliance on labels or approval, and expose structural dependencies that less explicit theories leave hidden. Empirical identification can strengthen, refine, restrict, or undermine applications of that architecture; it does not retroactively create the conceptual distinctions the architecture makes available.

Provenance, support graphs, implicit subject bindings, and challenge admission are latent functional relations. Verbal content, overt choice, and self-report may underdetermine them. Self-report is one fallible indicator, not privileged access. Empirical classification should therefore rely on converging interventions, transfer patterns, process measures, and model comparison. Persistent inability to identify the central provenance distinction would limit the theory's scientific utility even if its formal architecture remained coherent.

8.1 CET, control recruitment, encoded existence, and non-negligibility

Test control recruitment independently of the target output or behavior. Pre-specify the control-state indicators, aggregation H_ctrl , intervention family, operating measure μ^* , outcome map g , tolerance ϵ , comparison class, and sampling frame. At least one control-state indicator must track a standing policy state, action model, or control disposition without being defined by the same action, choice, or judgment later used as the predicted consequence. Prefer pre-output process measures, internal-state probes, transfer to held-out tasks, or interventions whose effects can be assessed on a separately measured control variable. When indicators disagree, the declared weighted, Pareto, or lexicographic rule must determine the comparison before the target classification is observed.

H1 tests a semantic-to-control association; it is not a complete test of encodedness. H1b separately tests the defeasible bridge from active, control-recruited correct-or-relevant representation to an outcome- evaluative realization. Formal EncOutcomeEval classification, actual existence (H15a), and non-negligible occurrence (H15b) are different claims. H15a and H15b use the same representation-token population and sampling unit; a person-level existential case is not a prevalence

result. Sampling, measurement, uncertainty, and the meaning of non-negligible remain to be governed prospectively rather than supplied as a universal number here.

H14 tests only the registered association between semantic differentiation and control-sensitive organization. An evolutionary or learning explanation is optional rationale, not a prediction of H14 and not evidence for pervasiveness, adaptation, or encoded CET. Objective welfare, analyst preference, and post-hoc behavioral success cannot substitute for the system's independently measured control dispositions.

8.2 Structural candidates, actual contribution, and empirical resolution

Identify the token event, frame activation, decision model, active routes, imports, and candidate structure independently. A candidate should fail to output when every actually contributing chain is blocked, when a route is not provenance-closed, or when anchoring fails. Split relationship, subject, and object witnesses must not pass moral typing.

ActualContribution is constitutive in the model and is evaluated only inside an evaluable fixed model instance. Empirical work pre-specifies a ContributionIDModel $M = \langle I_test, \Phi_obs, C_test \rangle$ and one or more candidate realization-profile schemes for a declared comparison class. Methodological admissibility asks whether a scheme was fixed independently and applied across its declared domain. Constitutive adequacy asks whether it preserves independently specified functional organization. Neither property is inferred from the desired moral classification.

Any future contribution scheme is fixed prospectively for its declared comparison class. Schemes may differ across genuinely distinct classes, but not across individual verdicts merely because the verdicts differ. The empirical contents of those schemes remain open.

Redundant support requires more than simple but-for dependence. Evidence may include activation during the formation interval, membership in a minimal sufficient active set, and changes in latency, confidence, affective intensity, persistence, appraisal structure, or other dimensions independently tied to functional organization even when sentence content is preserved. An active route that is profile-silent under the fixed adequate scheme is noncontributing and cannot moralize the token.

EmpiricallyUnresolved($\kappa, \{H_1, \dots, H_n\}, E$) when evidence E fails to identify one live provenance model under the pre-specified ContributionIDModel

Empirically unresolved is a legitimate outcome. Rival methodologically admissible schemes that preserve different candidate organizations and disagree about output-relevant contribution require discriminating evidence. Until an organization-preserving scheme or contribution-equivalent class is identified, the application is unresolved rather than negative. Invalid model instances are likewise not token-level nonoutputs.

8.3 Unified constraints, priority graphs, override, and tragedy

Test hard constraints, licensed overrides, residue, and tragic fallback separately. When Acc_delta is empty, manipulate the represented priority and violation profile while holding outcome scores fixed. Persistent outcome-only selection would falsify the violation-sensitive fallback family but not the general NDT architecture.

8.4 Appraisal targets and the represented-subject condition

The generic grammar predicts positive and negative appraisals of many target sorts. Moral typing adds a stricter claim: every moral token contains a represented moral subject, explicit in the content or implicit in the actual derivation. A natural disaster may be appraised as terrible without moral type; the same event becomes morally appraisable only when a person, collective, institution, or other represented agent is linked to it.

A direct empirical pressure test uses a subject model fixed independently of the target verdict. A robust judgment episode for which validated measures provide sufficiently strong negative evidence for both explicit and implicit subject representation under a declared detection model counts as adverse evidence against the subject condition when independent process, transfer, and contrast measures place the episode with paradigm moral judgments rather than tragic or axiological appraisals. The detection model must state its temporal window, unit, domain, sensitivity, false-negative risk, and evidential or equivalence threshold. The result is not automatically decisive: measurement error, hidden subject representations, and competing nonmoral classifications remain live possibilities. A fully specified superior rival strengthens the challenge but is not required before the evidence becomes adverse. No lexical verdict settles the classification by itself. Institutional cases must distinguish represented collective agency from grammatical personification.

8.5 MDT and the relationship-provenance identification program

Stage A fixes the builder, access, route, identity-form, and subject-gate measurement models before the target judgment. Positive and negative inherent value, perceived agency, non-scalar Relational Positioning, functional access, and attributed moral vision may not be inferred from moral language, one another, or downstream requirements. The inherent-valence profile must remain two-component and category-preserving throughout. Stage B asks whether a qualifying relationship actually helped construct or sustain the requirement, ordering, acceptability boundary, or appraisal used.

The central identification program compares three pathways while holding content and overt response constant: (1) a locally relationship-constructed requirement $m_{x,o} \rightarrow u_{x,o}$; (2) a general norm retaining relationship provenance $m_{class} \rightarrow u_{general} \rightarrow u_{x,o}$; and (3) provenance-free rule application $u_{policy} \rightarrow u_{x,o}$.

- Intervention: manipulate relationship representations while the explicit rule remains available.
- Transfer: compare generalization by relational similarity with transfer by formal rule membership.
- Process: compare latency, confidence, attention, affective intensity, persistence, memory, and sensitivity to reframing.
- Lesion and rescue: weaken the relationship route while preserving the norm cue, and vice versa.
- Model comparison: test relationship-provenance, policy-only, outcome-only, and mixed-route models on held-out data.

No one measure is assumed to identify provenance. If the pathways remain indistinguishable under every feasible intervention and process measure, MDT's empirical usefulness is correspondingly limited.

A bounded alternative interpretation applies to some third-person and first-person reconstructive cases: a participant may be assessing an agent's character or an action's moral worth rather than classifying the cognitive provenance of the focal judgment. That possibility belongs in the rival-model design, not in the constitutive definition. First-person prospective cases - where a person is deciding what to do before acting - remain the discriminating control; independent measures of appraisal target should separate the interpretations in reconstructive and third-person cases. NDT can model character assessment as agent-directed appraisal, but it does not yet formalize a probabilistic relation among character attribution, expected conduct, behavior evidence, and judgment type.

Relation-qualifier identification. A represented ownership or other qualifier should not be inferred merely from legal status, possessive wording, control of access, or intimacy. Identify the represented relation independently and test which downstream locus it affects. Candidate contrasts include stewardship toward pets versus otherwise similar wild animals, guardianship versus property, self-directed bodily sovereignty, institutional custody, and dominion-oriented frames. These contrasts are not assumed to be morally or psychologically equivalent. The target claim is that a represented relation can selectively change duties, permissions, priority, applicability, appraisal target, action description, or another registered downstream locus while the core relationship and overt content are held fixed. A change to an agent, object, or other core-individuating construal is core-changing for that operation and cannot satisfy StrictBender_s.

8.6 Moral subjecthood, principles, and third-personal judgment

The crossed high/low attributed-recognition traces provide the principal third-personal contrast. Hold the referent label, object frame, action, and core builder profile independently variable. Human and nonhuman referents should both clear the subject gate under high attributed integrative recognition and both fail it under low attribution. Answerability and responsibility are separate conditions for blame, not proxies for entity kind.

A sophisticated universalization or welfare calculation is not automatically moral. It may retain live relationship provenance, or it may operate as a free-standing formal or institutional routine. Philosophical self-description does not decide the issue. Conversely, the theory does not predict how common relationship-free principle use is; prevalence is empirical.

Collective moral-subjecthood hypothesis. Holding policy content, harm, and member conduct fixed, greater attributed collective recognition - the degree to which the institution is represented as able to register reasons, integrate standing, and revise conduct - should increase institution-directed blame relative to blame directed only at individual members. Weakening that attribution should shift blame toward members or dissolve institutional moral subjecthood.

H_collective: higher $MV_s(x_{institution,o,t,f,c})$, under fixed conduct and harm, predicts a higher institution-directed / member-directed blame ratio.

8.7 Case anchoring, alignment, plurality, and mode-indexed polyphony

Identify frame and token events independently. Polyphony requires distinct activations, a relevant frame difference, common case anchoring, the necessary beta, gamma, tau, or shared-Q bridge, local coherence, and conflict. Manipulating the bridge while holding local outputs fixed should move a pair between unaligned plurality and represented polyphony.

8.8 Challenge admission, arbitration, and practical limits

Admission should be measured before the target output through recognition or recall, target identification, credibility, and challenge-relevance measures, supplemented where possible by interference, attention, memory, or neural signatures. Output change must not define admission. Only after admission is independently established should skeptical defeat predict suspension or route failure.

Where subjects produce a global verdict, the model predicts evidence of a higher-order arbitration frame. Where none is represented, unresolved polyphony remains possible. Some token-level provenance questions may remain practically intractable; the article should set expectations accordingly rather than promise effortless classification.

8.9 Architecture-level instantiation and provenance-based kind unity

Full-architecture instantiation (H17). NDT's descriptive application includes the hypothesis that at least some real judgment episodes instantiate the complete framed-detachment architecture rather than only isolated modules. An identified case requires an independently identified frame, evaluative structure, decision or appraisal model, candidate token, and actually contributing live chain that jointly produce the registered output. A moral case additionally requires the represented relationship, subject binding, object binding, and same-chain local-contribution conditions. Component-wise compatibility alone does not establish full-architecture instantiation.

Provenance-based kind unity (H18). When token groups are identified independently of the target outcome, NDT predicts that provenance-based partitions will show greater held-out intervention, transfer, process, or explanatory stability than partitions based only on content, lexical labels, overt behavior, character appraisal, or free-standing policy source. Repeated failure across feasible pre-specified tests - especially consistent outperformance by rival partitions - counts against the claim that provenance supplies the proposed unifier. Causal or process heterogeneity by itself is not decisive, because one functional kind may be multiply realized.

9. Open Items

- Empirical identification of the control profile, aggregation rule, and boundary between encoded and attributed evaluation.
- Measurement models for positive and negative inherent value, perceived agency, functional access, positioning route and identity-form uptake, and moral vision, including abstract and collective objects.
- Actual contribution under redundant support, including permissible intervention contingencies and fine-grained token profiles.
- Reliable identification of explicit and implicit moral subjects in linguistic, behavioral, neural, and reconstructed data.
- Alternative fallback orderings for tragic choice.
- The full generic appraisal ontology and its relation to specialized types.
- Frame and judgment-event anchoring, action/target/dimension alignment, and polyphony in empirical data.
- Admission policies and alternative challenge semantics.
- The boundary between empirically unresolved provenance and evidence against the proposed distinctions.
- Reliable measurement of attributed collective recognition and its relation to institution- versus member-directed blame.
- Probabilistic links among character assessment, expected conduct, behavior evidence, and first-person or third-person judgment; this is an open extension, not an adopted current module.
- Executable and proof-assistant coverage of the full quantified theory.

- Noncircular identification of organization-preserving realization-profile schemes, contribution-equivalence conditions, and principled empirical unresolvedness.
- Identification of relation qualifiers and their operative locus, including stewardship, sovereignty, dominion, mixed modes, and cases where ownership causally changes a core builder rather than only downstream entailments.
- Independent demonstration that at least some real judgment episodes instantiate the complete framed-detachment architecture rather than isolated compatible components.
- Held-out comparison of provenance-based token partitions with lexical, behavioral, character-based, content-based, and policy-source partitions, including registered conditions under which the provenance unifier fails to earn its explanatory role.

10. Relation to Existing Work

NDT does not claim that its ingredients are individually unprecedented. Its current comparative conclusions are limited to the construct-level relations governed by NDT in Context Dataset 1.20. A named source, a bibliographic citation, or a registered source passage does not by itself establish a comparison relation. The deferred notes below identify questions for later adjudication; they are not active classifications, inheritances, novelty conclusions, or claims about what the cited work establishes.

10.1 Grounded argumentation and deferred defeat comparisons

Dataset 1.20 record CTX005 classifies Dung's grounded argumentation semantics, as bounded by Source Register 2.15 record S05, as an exact submodule of NDT's grounded challenge semantics. The least-fixed-point calculation is exact for the admitted attack graph. NDT separately specifies challenge admission, target typing, provenance, and moral classification; those additional features are outside the exact-submodule relation. Comparisons with Pollock's defeat terminology and with other structured-argumentation systems remain deferred candidates in Comparison Candidate Queue 1.1. No current Dataset relation is assigned to them.

10.2 Deferred comparison: defaults, priorities, and residue

The relationship among NDT's priority, override, and residue machinery and Horty's defaults-and-reasons framework remains a deferred candidate in Comparison Candidate Queue 1.1. Source Register 2.15 record S43 identifies one source passage, but no current Dataset record adjudicates a construct-level relation. No active comparison or inheritance claim is made here.

10.3 Deferred comparison: Universal Moral Grammar

The relationship between Universal Moral Grammar and NDT's typed judgment grammar or decision-model machinery remains a deferred candidate in Comparison Candidate Queue 1.1. No claim-specific Source Register unit and no current Dataset relation govern that comparison. The present specification therefore makes no predecessor, implementation, limitation, ecological- coverage, or explanatory-priority claim about Mikhail's program.

10.4 Deferred comparison: model-free, model-based, and Pavlovian accounts

Comparisons between NDT and the model-free, model-based, or Pavlovian accounts discussed by Cushman and Crockett remain deferred candidates in Comparison Candidate Queue 1.1. Dataset 1.20 record CTX051 concerns the distinct, source-bounded Cushman and Young (2011) construct; it does not govern the 2013 algorithmic taxonomy. No current implementation-family or rival-definition relation is assigned here.

10.5 Bayesian moral learning: Kleiman-Weiner, Saxe, and Tenenbaum

Dataset 1.20 record CTX016 classifies the hierarchical Bayesian moral-weight-inference construct bounded by Source Register 2.15 record S21 as an implementation family for NDT's learned social-value weighting. The learned weights can populate an outcome ordering while NDT's token-level moral-type test remains separate. This bounded relation does not classify the source as a whole theory, establish empirical adequacy, or imply that a learning algorithm determines moral type.

10.6 Deferred comparison: norm-expressivism and plans

The relationship between Gibbard's norm-expressivism and NDT's descriptive token-formation architecture remains a deferred candidate in Comparison Candidate Queue 1.1. No current Dataset relation governs that comparison. No active same-level contrast, explanatory limitation, or provenance conclusion is made here.

10.7 Deontic detachment

The name detachment echoes the logical problem of when conditional norms and facts license an unconditional obligation. NDT is not a new dyadic deontic logic; within this specification, it names a cognitive formation event under represented evaluation, feasibility, constraints, provenance, and defeat. Comparative claims involving factual, deontic, defeasible, or argumentation-based detachment remain deferred through Comparison Candidate Queue 1.1. No current Dataset relation or literature-derived inheritance claim is assigned in this subsection.

10.8 Existing moral-psychological theories

Moral Foundations Theory, Theory of Dyadic Morality, sentimental-rules accounts, Social Domain Theory, morality-as-cooperation, person-centered judgment, dual-process accounts, and moral-circle research each capture real regularities. In the source-backed formulations represented by the governed comparison corpus, no exact match was found for NDT's proposed token-level, provenance-based constitutive criterion. This is a bounded comparison, not a demonstration that no version or extension of those families could supply such a criterion. Values and themes do not determine their objects; dyadic accounts can represent self-directed, abstract, or apparently victimless cases by construing patients and harms, but the template still does not determine token-level moral type or distinguish moral from nonmoral appraisals with the same dyadic content; affect-linked prohibitions may implement or contribute to a judgment without fixing its type; cooperation is neither necessary nor sufficient; signature judgments can be scripted or heteronomous; and process labels do not identify provenance. NDT absorbs useful content, signature, learning, affective, and implementation modules while rejecting their elevation into exhaustive definitions.

Longstanding religious and philosophical traditions distinguish outward conformity from the source or spirit of conduct; Kant's shopkeeper is one familiar illustration. NDT neither derives its provenance criterion from that example nor adopts its concern with moral worth: actual contribution classifies judgment tokens descriptively rather than assessing an action's worth or an agent's purity of motive.

Executable status. Executable Partial Reference Model 2.2 is the finite companion for this revision. It preserves the Model 2.1 architecture while adding the subject-witness and counterpart-relationship-contributions required by Revision 1.6.3. It must preserve the positive and negative valence components through every gate and contribution test; distinguish the four profile states; exercise the identity-mediated, intimacy, and direct encounter route families plus all three identity forms; keep access separate; test exact complementary-role binding; and separate strict downstream bending from core-changing realization, input, or modulation. Those Gate 2 controls supplement rather than erase the Gate 1 contribution, binding, system, definedness, override, and nontriviality seams. The executable program does not discover latent routes from raw data, prove the quantified theory, or establish empirical identifiability or prose-code equivalence.

Appendix A. Compact Formal Summary

This appendix collects the core definitions without the surrounding defense. Subject index s is suppressed only where no ambiguity results.

A.1 Semantic, outcome, and practical evaluation

$\text{SemTEStructure}(r,c) := \langle \text{Sigma}_r, \text{sem}_r \rangle$; $\text{SemTES}(r,c) \leftrightarrow r$ instantiates that structure in c

$\text{OutcomeEval}(e,c) := \langle \text{Omega}_c, \text{Delta}(\text{Omega}_c), \text{out}_e, \text{pros}_e, \text{c}, \text{ReqOut}_e \rangle$, with the prospect order agreeing with the outcome order on degenerate prospects and no universal expected-utility lift

$\text{CtrlProfile}_s(D1,c) := \langle \text{Select}, \text{Maintain}, \text{Approach}, \text{Restore}, \text{Stabilize} \rangle$

$D1 \text{ } \triangleright \text{ctrl}_s, c \text{ } D2 \leftrightarrow H_{\text{ctrl}}(\text{CtrlProfile}_s(D1,c)) \triangleright H_{\text{ctrl}}(\text{CtrlProfile}_s(D2,c))$, with H_{ctrl} fixed before testing

$\text{StakeMap}_{\text{epsilon}}(g,r,c|\mu^*) \leftrightarrow \text{ArchitecturallyInstantiated}(g,r,c)$ and $\text{Pr}_{\mu^*}[g(\text{sigma1}) \triangleright \text{ctrl } g(\text{sigma2})] \triangleright 1 - \text{epsilon}$ and $\text{Pr}_{\mu^*}[g(\text{sigma1}) > \text{ctrl } g(\text{sigma2})] > 0$

$\text{EncOutcomeEval}(r,e,c) \leftrightarrow \text{OutcomeEval}(e,c)$ and exists $g[\text{StakeMap}_{\text{epsilon}}$ and $\text{Realizes}(e,g,r,c)$ and $\text{ConstitutiveConsumerRole}(g,e,r,c)$

Realizes and ConstitutiveConsumerRole are the typed realization and constitutive-consumer relations defined in §2.3. Universal strong semantic CET is an acknowledged external argument, not an NDT assumption. EncodedExistence and EncodedNonNegligible are separate empirical claims.

PrimitiveEvalBase(c):={e:exists r EncOutcomeEval(r,e,c) or exists r AttrOutcomeEval(r,e,c) or ExplicitEval(e,c)}

EvalClosure(f,c):=Cl_A({e in PrimitiveEvalBase(c):AccessibleTo(e,f,c)}), with Prov(e) defined

A.2 Decision models, perspectives, object scope, constraints, priority, and guidance

FrameOf(δ)=f; AgentOf(δ)=x; AgentReferentOf(δ)=y_A; TimeAnchorOf(δ)=t; OrdOf(δ)=e; PerspectiveOf(δ)= π

SupportGraphOf(δ ,c)= Γ_δ ; FallbackOf(δ) is an optional represented preorder φ_δ over Feas_ δ ; Basis(δ) and Warrants(δ) are nodes/links in Γ_δ

Req^{*}_Δ:Req^{*}act_Δout_Δ; h_Δ:Feas_Δ→Δ(Ω_Δ); $a \geq b \leftrightarrow h_\Delta(a) \geq_{\text{pros}_{e,c}} h_\Delta(b)$

ObjectScope_s(δ ,c):={o,y>: outcome, action, or requirement in δ is represented as concerning o/y}

Violates(a,u, δ ,c)↔ApplicableRequirement_s(u,a, δ ,c) and SatisfactionDeterminate_ δ (a,u,c) and not Satisfies_ δ (a,u,c); missing or unresolved satisfaction does not count as false

PriorityWellFormed(δ)↔finite∧irreflexive∧acyclic; LicensedOverride, Disqualifies, StructuralResidue, and token-relative OutputResidue as in §3.2

Acc_ δ :={a∈Feas_Δ:¬∃u Disqualifies(u,a, δ ,c)}

Guide_ δ :Max_Δ≥ δ (Acc_ δ) if Acc_ δ ≠∅; otherwise Max_Δ≥fallback δ (Feas_Δ) when φ_δ is represented

TragicChoice(a, δ)↔a∈Guide_ δ ∧a∉Acc_ δ ; StructuralResidue records model-level license or tragedy, while OutputResidue requires operative route support or tragedy in an output token

A.3 Appraisal, structural formation, modes, and candidates

j::=OughtChoose|Violation|Wrong|Blameworthy|Praiseworthy|Appraise(z_A,q,pol,t)

TargetSort(z_A) in {agent,action,agent-act,outcome,event,object,relation,institution,way-of-life}; AgentActTarget(x,a) has sort agent-act

AppStruct_delta(q):=<X_delta,q,>=app_delta,q.PosRegion_delta,q.NegRegion_delta,q.Scope,Target,Prov>

AssessmentTargetFormed_s(c,kappa,delta,z_A,t) requires target-sort match, route-carried target binding, case/frame anchoring, and TimeAnchorOf(delta)=t

PraiseCandidate uses Profile_delta,q(AgentActTarget(x,a),t,c), ResponsibleCapacityFor_s on q, and CreditsFor_s; favorable appraisal without attributed credit remains generic Appraise; GenericAppraisalCandidate uses Profile_delta,q(z_A,t,c)

Candidate ↔ Prospective or Violation or Wrong or Blame or Praise or GenericAppraisal

Generic appraisal does not imply moral type. Every moral token separately requires a represented subject explicit or implicit in its actual derivation.

A.4 Frame-relative agent and object construals, moral vision, and represented moral relationships

AgentConstrual_s(x,y_A,f,c): frame f presents agent referent y_A under agent construal x

AgentConstrual_s(x_self,y_self,f,c) and ObjectConstrual_s(o_self,y_self,f,c) is role co-reference; role-correct subject and object bindings remain distinct

AgentConstrual_s(x,y_A,f,c)∧AgentConstrual_s(x,y'_A,f,c)→y_A=y'_A;

AgentAnchorOf_s(x,f,c)=y_A↔AgentConstrual_s(x,y_A,f,c)

CapableAs_s(y_A,x,Φ,c): s attributes to agent referent y_A, under construal x, the capacity to instantiate profile Φ under suitable conditions

ObjectConstrual_s(o,y,f,c): frame f presents object referent y under object construal o

IRAttributed_s(y_A,x,o,t,f,c')↔RepresentsInherentValenceAs_s(y_A,x,o,t,f,c')∧RepresentsRelevantRelationsAs_s(y_A,x,o,t,f,c')

$\wedge \text{RepresentsAgencyModelAs}_s(y_A, x, o, t, f, c') \wedge \text{IntegratesToEvalStructureAs}_s(y_A, x, o, t, f, c')$
 $MV_s(x, o, t, f, c) := \text{Degree}_s[\text{CapableAs}_s(\text{AgentAnchorOf}_s(x, f, c), x, \lambda c'. \text{IRA} \text{Attributed}_s(\text{AgentAnchorOf}_s(x, f, c), x, o, t, f, c'), c)]$
 $\text{InherentValenceProfile}_s(o, t, f, c) := \langle \text{PositiveInherentValue}_s(o, t, f, c), \text{NegativeInherentValue}_s(o, t, f, c) \rangle$
 $\text{PositiveOnlyValence}_s \leftrightarrow \text{PIV} > 0 \wedge \text{NIV} = 0$; $\text{NegativeOnlyValence}_s \leftrightarrow \text{PIV} = 0 \wedge \text{NIV} > 0$; $\text{MixedValence}_s \leftrightarrow \text{PIV} > 0 \wedge \text{NIV} > 0$;
 $\text{NeitherValence}_s \leftrightarrow \text{PIV} = 0 \wedge \text{NIV} = 0$
 $\mathcal{M}_R := \langle \mathcal{U}_{\text{plus}}, \mathcal{U}_{\text{minus}}, \theta A, \mathcal{J}_{\text{plus}}, \mathcal{J}_{\text{minus}}, \mathcal{G}_{\text{plus}}, \mathcal{G}_{\text{minus}}, \mathcal{R} \rangle$; $\mathcal{M}_S := \langle \theta M \rangle$
 $\text{ValenceGate}_s(o, t, f, K, c) \leftrightarrow \text{ValenceBranchesAdmissible}_s(\text{PositiveValenceCriterionOf}_s(o, t, f, c), \text{NegativeValenceCriterionOf}_s(o, t, f, c), K, c)$
 $\wedge (\text{PositiveValenceBranchPass}_s(o, t, f, c) \vee \text{NegativeValenceBranchPass}_s(o, t, f, c))$
 $\text{CoreGate}_s(x, o, t, f, K, c) \leftrightarrow \text{ValenceGate}_s(o, t, f, K, c)$
 $\wedge \text{AgencyCap}_s(x, o, t, f, c) \geq \theta A$
 $\text{PositiveRelationshipStrength}_s(x, o, t, f, c) := \text{PositiveAgencyInteractionOf}_s(x, o, t, f, c)$
 $(\text{PositiveInherentValue}_s(o, t, f, c), \text{AgencyCap}_s(x, o, t, f, c))$
 $\text{NegativeRelationshipStrength}_s(x, o, t, f, c) := \text{NegativeAgencyInteractionOf}_s(x, o, t, f, c)$
 $(\text{NegativeInherentValue}_s(o, t, f, c), \text{AgencyCap}_s(x, o, t, f, c))$
 $\text{RelationshipStrengthProfile}_s(x, o, t, f, c) := \langle \text{PositiveRelationshipStrength}_s(x, o, t, f, c), \text{NegativeRelationshipStrength}_s(x, o, t, f, c) \rangle$
 $\text{RelationshipStrengthGate}_s(x, o, t, f, K, c) \leftrightarrow \text{RelationshipStrengthBranchesAdmissible}_s(x, o, t, f, K, c)$
 $\wedge (\text{PositiveRelationshipStrengthBranchPass}_s(x, o, t, f, c) \vee \text{NegativeRelationshipStrengthBranchPass}_s(x, o, t, f, c))$
 $\text{CoreMRCondition}_s(x, o, t, f, K, c) \leftrightarrow \text{CoreGate}_s(x, o, t, f, K, c) \wedge \text{RelationshipStrengthGate}_s(x, o, t, f, K, c)$
 $\text{RelationalPositioning}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c) \leftrightarrow \text{IdentityMediatedRP}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c)$
 $\vee \text{IntimacyRP}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c) \vee \text{DirectEncounterRP}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c)$
 $\text{RelationalPositioningGate}_s(m, x, y_A, o, y, t, f, c) \leftrightarrow \exists \lambda \text{bda} \text{ RelationalPositioning}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c)$
 $\text{RelationalPositioningOn}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, \text{chi}, c) \leftrightarrow \text{RelationalPositioning}_s(\lambda \text{bda}, m, x, y_A, o, y, t, f, c) \wedge$
 $\text{RelationalSubrouteOf}(\lambda \text{bda}, \text{chi}, c) \wedge m \in \text{Nodes}(\lambda \text{bda})$
 $\text{BuilderQualifiedMRCondition}_s(m, x, y_A, o, y, t, f, K, c) \leftrightarrow$
 $\text{CoreMRCondition}_s(x, o, t, f, K, c) \wedge \text{RelationalPositioningGate}_s(m, x, y_A, o, y, t, f, c)$

The inherent-valence and relationship-strength profile pairs remain present after their Boolean gates. No gate or relationship aggregation may replace either ordered pair with max, net, sum, mean, or another scalar. The relationship-strength profile is derived from the first two builders and is not a fourth builder.

$\text{MoralSubjectGate}_s(x, y_A, o, t, f, c) \leftrightarrow \text{SelfReferent}_s(y_A, c) \vee MV_s(x, o, t, f, c) \geq \theta M$

SelfReferent clears only the subject gate; the ordinary builder, token, binding, output, and same-chain provenance conditions still apply.

$\text{MoralRelationRep}_s(m, x, y_A, o, y, t, f, K, c) \leftrightarrow \text{RelationshipRep}_s(m, x, y_A, o, y, t, f, c)$

$\wedge \text{AgentConstrual}_s(x, y_A, f, c) \wedge \text{ObjectConstrual}_s(o, y, f, c)$

$\wedge \text{BuilderQualifiedMRCondition}_s(m, x, y_A, o, y, t, f, K, c)$

$\wedge \text{MoralSubjectGate}_s(x, y_A, o, t, f, c)$

$\text{MR}_s(x, y_A, o, y, t, f, K, c) \leftrightarrow \exists m \text{ MoralRelationRep}_s(m, x, y_A, o, y, t, f, K, c)$ [extensional shorthand]

FunctionalAccess_s(u0, psi, c) identifies whether support unit u0 can enter active process psi under a fixed access model.

IdentityMediatedRP_s, IntimacyRP_s, and DirectEncounterRP_s are the governed positioning route families.

SharedIdentityForm_s, ContrastiveIdentityForm_s, and ComplementaryRoleIdentityForm_s exhaust the governed identity forms. Functional access is not a positioning subtype. Standing-mediated is retired; an operative impersonal directed assignment may qualify through ComplementaryRoleIdentityRP_s. Contrastive identity does not imply negative inherent value, and downstream identity consequences require separate bender credits.

A philosophical label does not satisfy MoralRelationRep. Abstract objects are permitted only under independently identified construal, three-builder, and provenance conditions.

A.5 Connected route survival, import/export, and moral provenance

SupportGraphOf(z, c) = Γ_z ; Route(ρ, z, c) $\leftrightarrow$ ρ is a provenance-closed, derivationally sufficient, minimally sufficient subgraph of Γ_z

FunctionalOrganizationOf(κ, K, c) := $\langle$ active processes, dependency structure, transition/disposition profile, invariance conditions $\rangle$

MethodologicallyAdmissible(S, K, c) $\leftrightarrow$ fixed independently and classification-independent and within declared invariance domain

ConstitutivelyAdequate(S, K, c) $\leftrightarrow$ OrganizationPreserving(S, K, c), assessed prior to ActualContribution

Conditional convergence holds for adequate schemes connected over a nonempty registered intervention basis by a same-resolution profile bridge preserving profile structure, intervention effects, minimal sufficiency, and the defined chain-profile difference-maker verdict

CrossSchemeContributionInvariant is a substantive model-class constraint; it is not derivable from adequacy alone

ConstitutiveTokenProfile(κ, c) := RealizationProfile_ $\{S^* \{s, K^* \{\kappa, c\}\}(\kappa, K^* \{\kappa, c\}, c)$

ActualContribution(χ, κ, c) $\leftrightarrow$ ContributionVerdict_ $\{S^* \{s, K^* \{\kappa, c\}\}(\chi, \kappa, K^* \{\kappa, c\}, c)$

AccessFor_ $s(u_0, \chi, \kappa, c) \leftrightarrow$ exists ψ [ProcessCarriesOn($\psi, u_0, \chi, \kappa, c$) and FunctionalAccess_ $s(u_0, \psi, c)$]

LocalContributionVerdict_ $\{S^* \{s, K^* \{\kappa, c\}\}$ requires AccessFor_ s , reachability, local minimal-sufficient-set membership, and an evaluation-profile difference maker under $K^* \{\kappa, c\}$

ProfileSilent(χ, κ, c) $\rightarrow$ not ActualContribution(χ, κ, c)

FunctionallyActiveChain(χ, κ, c) $\leftrightarrow$ the frame, model, and token processes composing χ were functionally active and correctly anchored, independently of challenge survival

LiveChain(χ, κ, c) $\rightarrow$ FunctionallyActiveChain(χ, κ, c)

RealizedLiveChain(χ, κ, c) $\leftrightarrow$ LiveChain(χ, κ, c) and ActualContribution(χ, κ, c)

Outputs(c, κ) $\leftrightarrow$ EvaluableFor_ $s(\kappa, c)$ and Candidate(c, κ) and exists χ RealizedLiveChain(χ, κ, c)

InvalidModelInstance and EmpiricallyUnresolvedScheme withhold Boolean application; neither implies Outputs=false

ContributionIDModel(M) := $\langle I_test, \Phi_obs, C_test \rangle$; M and IndicatorMap supply evidence but do not constitute the target

BoundRelationalPositioning_ $s(m, x, y_A, o, y, t, f, c) := \{\lambda: RelationalPositioning_s(\lambda, m, x, y_A, o, y, t, f, c)\}$

BoundBuilderProfile_ $s(m, x, y_A, o, y, t, f, c) := \langle InherentValenceProfile_s(o, t, f, c), AgencyCap_s(x, o, t, f, c),$

BoundRelationalPositioning_ $s(m, x, y_A, o, y, t, f, c) \rangle$

RelationshipProfileContributes_ s requires MoralRelationRep and a RelationalPositioningOn witness for the exact tuple, tier-one local actual contribution by m to RelevantEval on the same chain, and, at tier two, either a controlled same-chain effect on the exact three-builder profile or a profile-level NESS witness for that profile. Profile-level NESS can preserve redundant profile support but cannot rescue a false tier-one local difference-maker verdict. No positioning route, identity form, or individual builder is required to be an individual but-for cause.

MoralSubjectOf_ $s(\kappa, x, y_A, \chi, c) \leftrightarrow$ explicit content-agent match or route-carried implicit subject-to-target link on χ

CarriesSubjectBinding_ s and SubjectTargetsOn_ s require the live subject binding and its path to the formed appraisal target in the same chain

In this précis, SubjectContributes_ s additionally names the exact typed subject-binding node's local contribution on that same live chain; a merely carried, mistyped, or profile-silent subject link does not qualify

IntegrationRoleFor(n_i, κ, c) requires an independently typed requirement, appraisal, acceptability, or other relation-bearing evaluation-input node; a bare guide, candidate, output, or post-hoc common descendant does not satisfy it

BoundProfileIntegrationOn_ $s(n_i, m, n_o, E, \chi, \kappa, c) \leftrightarrow$ IntegrationRoleFor(n_i, κ, c) and $\{m, n, n_o, n_i, E\}$ subseq Nodes(χ) and Reach_ $\chi(m, n_i, c)$ and Reach_ $\chi(n, n_i, c)$ and Reach_ $\chi(n_o, n_i, c)$ and Reach_ $\chi(n_i, E, c)$

This common-integration condition does not require a subject-binding or object-binding node to be an ancestor of the relationship node. It rejects only the looser case in which the three routes first meet at a generic guide, candidate, or output.

TimeInOrAnchors(j, δ, t) $\leftrightarrow$ TimeIn(j, t) or [no time occurs in j and TimeAnchorOf(δ) = t]

MoralJudgment(c, κ) $\leftrightarrow$ EvaluableFor_ $s(\kappa, c)$ and Outputs(c, κ) and exists

$K, \chi, m, x, y_A, o, y, t, f, n, n_o, n_i$ [K = ComparisonClassOf(κ, c) and MoralRelationRep_ $s(m, x, y_A, o, y, t, f, K, c)$ and

MoralSubjectOf and SubjectBindingNode_s(n,x,y_A,kappa,chi,c) and LocalActualContribution(n,RelevantEval(kappa),chi,kappa,c) and ObjectBindingNode_s(n_o,o,y,chi,BaseOf(kappa),c) and BoundProfileIntegrationOn_s(n_i,m,n_o,RelevantEval(kappa),chi,kappa,c) and RelationshipProfileContributes_s(m,x,y_A,o,y,t,f,RelevantEval(kappa),chi,kappa,K,c) and SubjectContributes_s and BearsOnOn and frame/time matches on one RealizedLiveChain]

The same node n is the subject-binding input to common integration and a local-actual-contribution witness; separate existential witnesses cannot be composed.

An Appraise token for which the system represents no agent as subject may output but cannot satisfy MoralSubjectOf and is therefore nonmoral

EmpiricallyUnresolved(kappa,H,E) when evidence E does not identify one support/provenance hypothesis

A.6 Grounded attack semantics and operative override

AdmittedChallenge_s(d,u,c) $\leftrightarrow$ Registered(d,c) and Targets(d,u,c) and Credibility(d,c) $\geq$ theta_cred and ChallengeRelevance(d,u,c) $\geq$ tau_chal

Grounded iteration over admitted meta-attacks yields IN, OUT, and UNDECIDED; Defeated(d,c) $\leftrightarrow$ d is OUT

PremiseAvailable and LinkAvailable require every admitted undercutter to be defeated; DirectClear treats target attacks likewise

LicensedOverride(u,a,delta,c) is structural: the formed model contains a qualifying violation, OverrideGround, and priority path licensing the override for action a in model delta. OperativeOverrideOn(u,a,v,w,chi,kappa,c) is token-relative and additionally requires a ground node v and warrant w on a surviving, actual contributing chain with a path to RelevantEval(kappa). A licensed override need not be operative on a given token route.

OutputResidue(a,u,kappa,c) $\leftrightarrow$ Outputs(c,kappa) and BaseOf(kappa)=delta and Violates(a,u,delta,c) and [exists v,w,chi OperativeOverrideOn(u,a,v,w,chi,kappa,c) or TragicChoice(a,delta)]

A.7 Case anchoring, alignment, plurality, conflict, and polyphony

SameJudgmentToken $\leftrightarrow$ same judgment-event anchor; SameFrameToken $\leftrightarrow$ same frame-event anchor; FrameTypeEquivalent is separate

SameCase_s(kappa1,kappa2,c) $\leftrightarrow$ exists zeta shared by both CaseOf projections, with each token's own frame and episode construal mapping to that zeta

Aligned(beta,delta1,delta2,c): surviving action-alignment route over shared agent and temporal anchors

AlignmentCovers(beta,kappa1,kappa2,c) $\leftrightarrow$ every action in MentionedActs(kappa1) and MentionedActs(kappa2) has a nonempty beta execution map

PairAligned(beta,kappa1,kappa2,c) $\leftrightarrow$ Aligned and AlignmentCovers; missing maps yield unaligned plurality, never conflict

TargetAligned(tau,kappa1,kappa2,c): surviving target-alignment route for the formed targets at BaseOf(kappa1) and BaseOf(kappa2)

DimAligned(gamma,q1,delta1,q2,delta2,c): surviving dimension-alignment route into one appraisal space

ProspectiveConflict binds each token's own ChoiceSetOf and BaseOf, requires PairAligned and nonempty execution extensions, and requires disjointness

GuideCondemnationConflict requires overlap with the condemned execution and exclusion from the condemning model's acceptable extension

Overall, aspectual, and generic-target appraisal conflict bind each token's question, dimension, target, polarity, and base projection and require the beta, tau, gamma, or shared-Q bridges appropriate to target sort. Generic-target opposition uses PolarityOf, not raw profile values.

PluralPair requires outputs, distinct frame/token events, common case, admissibility, and FrameDifferenceRelevant

ComparisonReady requires the bridge appropriate to pair type; UnalignedPair $\leftrightarrow$ PluralPair and not ComparisonReady

UnalignedPair subtypes include no action/target bridge and execution/target-aligned but appraisal-not-comparison-ready pairs lacking shared Q or surviving gamma

PolyphonicPair $\leftrightarrow$ PluralPair and local coherence and Conflict; Polyphony $\leftrightarrow$ exists pair

A.8 Composition witnesses

- A: first-person moral care; relationship, subject, and object paths occur in one actually contributing chain.
- B: same-content hygiene token outputs without moral provenance.
- C/D/Y/Z: crossed human/nonhuman referents under high/low attributed recognition; the subject-gate verdict tracks the represented profile rather than entity kind.
- AA/AB: access can block one relationship path without becoming a positioning subtype, while the two valence components and positioning remain independently represented.
- Positioning-deletion control: with the inherent-valence profile, perceived agency, access, subjecthood, and output held fixed, removing every Relational Positioning realization for the exact relationship makes RelationalPositioningGate and moral type fail; removing one witness need not do so when another route witness remains.
- Valence-profile controls: positive-only, negative-only, mixed, and neither remain distinguishable after the Boolean ValenceGate decision. Replacing the pair by net, maximum, sum, mean, or another scalar invalidates the model instance.
- Relationship-strength controls: positive-by-agency and negative-by-agency interaction branches remain an ordered pair after their Boolean gate. Replacing the pair by one scalar, omitting the load-bearing gate, or failing the frozen positive-only relationship-strength comparison invalidates the model instance.
- Identity controls: shared, contrastive, and complementary-role forms each require an independently specified operative schema and exact bridge. Removing a downstream duty does not by itself remove the identity bridge; removing the bridge does not follow from retaining the duty.
- Bender controls: a strict bender follows a live, accessible, actually contributing path through a typed downstream locus to the relevant evaluation. Its registered removal intervention preserves an equal exact before/after core snapshot. Residue is one of the twelve loci. A change to a builder or core construal is typed as realization, input, or modulation instead.
- E: agent-implicit institution or negligence appraisal carries a subject binding even when the sentence omits the agent.
- F: a disaster appraisal represented without any agent as subject may output as tragic or axiological but remains nonmoral.
- G: violation-sensitive tragic fallback selects by the represented violation profile before outcome tie-breaking.
- AD: moral self-relation carries role-distinct self subject and self object bindings on one actually contributing chain.
- AE/AG: invalid and empirically unresolved profile-scheme applications withhold classification rather than returning a negative token verdict.
- AF: an ownership qualifier bends downstream priority or override without directly conferring moral type.
- Countermodels: split witnesses, provenance stripping, unrelated chains, inactive or profile-silent moral backup routes, duplicate frame labels, unmapped-action pseudo-conflict, and principle-only moralization are rejected.

Bibliography

A. Work by the Author

Beal, B. (2018). *Raskolnikov's Confession: A Dostoevskian Model of Moral Psychology* (Doctoral dissertation, Emory University, Graduate Institute of the Liberal Arts).

Beal, B. (2020). What are the irreducible basic elements of morality? A critique of the debate over monism and pluralism in moral psychology. *Perspectives on Psychological Science*, 15(2), 273-290.

Beal, B. (2021). The nonmoral conditions of moral cognition. *Philosophical Psychology*, 34(8), 1097-1124.

Beal, B. (2022). The polyphony principle. *Behavioral and Brain Sciences*, 45, e222.

Beal, B. (2025, January 20). The Is-to-Ought model of moral cognition. Bree's Substack.

- Beal, B. (2025, January 31). On the relation between moral philosophy and moral psychology. Bree's Substack.
- Beal, B., & Gogia, G. (2021). Cognition in moral space: A minimal model. *Consciousness and Cognition*, 92, 103-134.
- Beal, B., & Rottman, J. (preprint). Ontological frames decisively outperform moral foundations in predicting moral judgments. SSRN abstract 4724806.

B. Philosophical Predecessors

- Anscombe, G. E. M. (1957/2000). *Intention*. Harvard University Press.
- Berlin, I. (2001). My intellectual path. In H. Hardy (Ed.), *The Power of Ideas*. Princeton University Press.
- Foot, P. (2001). *Natural Goodness*. Oxford University Press.
- Heidegger, M. (1927/2010). *Being and Time* (J. Stambaugh & D. Schmidt, Trans.). SUNY Press.
- Held, V. (2006). *The Ethics of Care*. Oxford University Press.
- Hume, D. (1739-1740). *A Treatise of Human Nature*, Book III, Part I, Section 1.
- Kant, I. (1785). *Groundwork for the Metaphysics of Morals*.
- Korsgaard, C. M. (1983). Two distinctions in goodness. *The Philosophical Review*, 92(2), 169-195.
- Millikan, R. G. (1984). *Language, Thought, and Other Biological Categories*. MIT Press.
- Noddings, N. (1984). *Caring: A Relational Approach to Ethics and Moral Education*. University of California Press.
- Putnam, H. (2002). *The Collapse of the Fact/Value Dichotomy and Other Essays*. Harvard University Press.
- Searle, J. R. (1964). How to derive "ought" from "is." *The Philosophical Review*, 73(1), 43-58.
- Searle, J. R. (1995). *The Construction of Social Reality*. Free Press.
- Thompson, M. (2008). *Life and Action*. Harvard University Press.
- Williams, B. (1985). *Ethics and the Limits of Philosophy*. Harvard University Press.

C. Empirical and Scientific Sources

- Nichols, S. (2002). Norms with feeling: Towards a psychological account of moral judgment. *Cognition*, 84(2), 221-236.
- Barrett, H., Bolyanatz, A., Crittenden, A., Fessler, D., Fitzpatrick, S., Gurven, M., et al. (2016). Intentions and moral judgment across societies. *Proceedings of the National Academy of Sciences*, 113(17), 4688-4693.
- Barrett, L. F. (2017). *How Emotions Are Made*. Houghton Mifflin Harcourt.
- Crimston, C. R., Bain, P. G., Hornsey, M. J., & Bastian, B. (2016). Moral expansiveness: Examining variability in the extension of the moral world. *Journal of Personality and Social Psychology*, 111(4), 636-653.
- Crimston, C. R., Hornsey, M. J., Bain, P. G., & Bastian, B. (2018). Toward a psychology of moral expansiveness. *Current Directions in Psychological Science*, 27(1), 14-19.
- Damasio, A. (1994). *Descartes' Error*. Putnam.
- Dung, P. M. (1995). On the acceptability of arguments and its fundamental role in nonmonotonic reasoning, logic programming, and n-person games. *Artificial Intelligence*, 77(2), 321-357.
- Friston, K. (2010). The free-energy principle: A unified brain theory? *Nature Reviews Neuroscience*, 11, 127-138.
- Gould, S. J., & Lewontin, R. C. (1979). The spandrels of San Marco and the Panglossian paradigm. *Proceedings of the Royal Society B*, 205(1161), 581-598.
- Haidt, J. (2001). The emotional dog and its rational tail: A social intuitionist approach to moral judgment. *Psychological Review*, 108(4), 814-834.
- Haslam, N. (2006). Dehumanization: An integrated review. *Personality and Social Psychology Review*, 10(3), 252-264.
- Haslam, N., & Loughnan, S. (2014). Dehumanization and inhumanization. *Annual Review of Psychology*, 65, 399-423.
- Law, K. F., Campbell, D., & Gaesser, B. (2022). Biased benevolence: The perceived morality of effective altruism across social distance. *Personality and Social Psychology Bulletin*, 48(3), 426-444.
- Marr, D. (1982/2010). *Vision*. MIT Press.

Much, N. C., & Shweder, R. A. (1978). Speaking of rules: The analysis of culture in breach. *New Directions for Child and Adolescent Development*, 2, 19-39.

Shweder, R. A. (1992). Ghostbusters in anthropology. In R. G. D'Andrade & C. Strauss (Eds.), *Human Motives and Cultural Models* (pp. 45-58). Cambridge University Press.

Tajfel, H., & Turner, J. C. (1979/2004). The social identity theory of intergroup behavior. In J. T. Jost & J. Sidanius (Eds.), *Political Psychology: Key Readings*. Psychology Press.

Tversky, A., & Kahneman, D. (1981). The framing of decisions and the psychology of choice. *Science*, 211(4481), 453-458.

D. Framing and Literary Sources

Bakhtin, M. M. (1929/1984). *Problems of Dostoevsky's Poetics* (C. Emerson, Ed. & Trans.). University of Minnesota Press.

Bermudez, J. L. (2022). Rational framing effects: A multidisciplinary case. *Behavioral and Brain Sciences*, 45, e220.

Dostoevsky, F. (1866/1993). *Crime and Punishment* (R. Pevear & L. Volokhonsky, Trans.). Vintage Classics.

E. Formal and Computational Neighbors

Alchourrón, C. E. (1996). Detachment and defeasibility in deontic logic. *Studia Logica*, 57, 5-18.

Beirlaen, M., & Straßer, C. (2018). Structured argumentation with prioritized conditional obligations and permissions. *Journal of Logic and Computation*, 29, 187-214.

Crockett, M. J. (2013). Models of morality. *Trends in Cognitive Sciences*, 17, 363-366.

Cushman, F. (2013). Action, outcome, and value: A dual-system framework for morality. *Personality and Social Psychology Review*, 17, 273-292.

Gibbard, A. (1990). *Wise choices, apt feelings: A theory of normative judgment*. Harvard University Press.

Horty, J. F. (2012). *Reasons as defaults*. Oxford University Press.

Kleiman-Weiner, M., Saxe, R., & Tenenbaum, J. B. (2017). Learning a commonsense moral theory. *Cognition*, 167, 107-123.

Mikhail, J. (2007). Universal moral grammar: Theory, evidence and the future. *Trends in Cognitive Sciences*, 11, 143-152.

Mikhail, J. (2011). *Elements of moral cognition*. Cambridge University Press.

Pollock, J. L. (1987). Defeasible reasoning. *Cognitive Science*, 11, 481-518.

Pollock, J. L. (1991). A theory of defeasible reasoning. *International Journal of Intelligent Systems*, 6, 33-54.

Strasser, C. (2011). A deontic logic framework allowing for factual detachment. *Journal of Applied Logic*, 9, 61-80.