

METHODOLOGY WHITE PAPER

Representing Contested Knowledge in a Public Reference System

The Divinity Atlas Methodology

CENTRAL PROPOSITION A reference work can represent what sources establish, what traditions hold, where accounts diverge, and what remains unknown without collapsing those categories into a single authoritative voice.

Chris Leite
Interactive Lion
Draft 1.1 | August 17, 2026

Prepared for external scholarly and methodological review

Creative Commons Attribution 4.0 International (CC BY 4.0)

Cite as: Divinity Atlas Methodology White Paper, Draft 1.1, August 17, 2026. Interactive Lion.
<https://www.divinityatlas.com/methodology>

Abstract

Divinity Atlas is a public reference system for religions, mythologies, philosophies, and esoteric traditions. Its central methodological problem is not merely how to store many claims, but how to represent claims that are differently established, differently held, disputed, culturally restricted, or genuinely unresolved. The system separates the evidentiary status of a claim from the stance a tradition takes toward it; treats sources, facts, relationships, uncertainty, and dissent as structured records; and subjects changes to auditable write procedures and executable quality gates. This paper describes that framework, its operational controls, and its public challenge and case process. It also states the framework's present limitations. Divinity Atlas has not yet undergone broad external academic validation, and internal second-pass review is not presented as equivalent to scholarly peer review. The system is offered as a transparent, testable methodology for public knowledge representation and as an invitation to scholars, librarians, and digital-humanities researchers to evaluate and improve it.

Document Status

This is a pre-publication methodology paper prepared from the active Divinity Atlas operational playbook, database constraints, quality-gate registry, audit trail, and public adjudication model. Counts are a launch-day snapshot taken August 17, 2026 and will continue to change. The paper distinguishes implemented controls from proposed external-review practices.

Contents

1. The problem: reference prose hides different kinds of certainty
2. Scope, principles, and unit of representation
3. Ontological architecture
4. Evidence status and tradition stance
5. Sources, citations, and reliability tiers
6. Unknowns, disagreement, and open questions
7. Editorial production and review workflow
8. Auditability, quality gates, and reproducibility
9. Cultural restrictions and rights-sensitive material
10. Public challenge, cases, corrections, and dissent
11. Automation, human judgment, and AI disclosure

12. Current scale and operational evidence

13. Limitations and validation agenda

14. Conclusion

Appendix A. Technical control map

Appendix B. Proposed external scholarly-review protocol

References

1. The Problem: Reference Prose Hides Different Kinds of Certainty

A conventional encyclopedia paragraph often combines several logically different acts: identifying an object, asserting a historical fact, summarizing a community's self-understanding, reconciling competing accounts, and deciding what level of confidence a reader should place in the result. When these acts are compressed into fluent prose, disagreement can appear to have been resolved merely because the prose has one grammatical voice.

Religion makes this problem especially visible. A statement can be historically well documented while rejected within a tradition. Two traditions can describe the same subject differently without either functioning as the default account. Scholars can agree that competing interpretations exist while disagreeing about which interpretation is strongest. A source can be useful evidence for what a community teaches without being independent evidence that the teaching is historically true.

Divinity Atlas responds by separating representation into structured layers. The system records the thing being discussed, the property or relationship being asserted, the source supporting the assertion, the evidentiary status assigned to it, and any tradition-specific stance. Public prose is generated from or attached to those records, but prose is not permitted to erase their distinctions.

DESIGN QUESTION How can a public reference work state what is known, show who holds what, preserve disagreement, and expose what would change its conclusions?

2. Scope, Principles, and Unit of Representation

The atlas covers religious traditions, denominations, deities, historical figures, sacred texts, doctrines, practices, rituals, places, artistic works, philosophical systems, divinatory systems, and related subjects. The model is intentionally extensible: a new kind of subject is registered as data rather than introduced through a new physical table and application-specific code path.

2.1 Governing principles

- Separate evidentiary status from community or tradition stance.
- Prefer explicit uncertainty to plausible completion.
- Require a traceable source for established public factual claims, subject to narrow documented exceptions for descriptive and structural records.
- Treat disagreement as data rather than as an editorial inconvenience.
- Preserve changes, rulings, dissents, and corrections instead of silently replacing them.
- Fail closed when rights, cultural restrictions, vocabulary, or provenance are unresolved.
- Make operational rules executable where reliable automation is possible, and name where judgment remains necessary.

- Distinguish internal review independence from external scholarly validation.

2.2 The claim atom

The smallest reviewable unit is a claim atom: one entity, one allowed property, one typed value, one evidentiary status, optional explanatory note, optional tradition context, and a source reference when required. A relationship is a parallel atom joining a subject entity to an object entity through a typed relation and carrying its own status, stance, note, and source.

This granularity enables a challenge to target a specific assertion rather than an entire page. It also permits a page to contain established, debated, debunked, and unresolved claims simultaneously without assigning one status to the entry as a whole.

3. Ontological Architecture

Everything described by the atlas is represented as an entity. An entity has a stable identifier, a display name, a stable route key, and an entity type. Each entity type declares the properties its members may hold and the native data type of each property. Values are stored in typed tables for text, integers, decimals, dates, booleans, and controlled enumerations.

Relation types are themselves entities. Their immutable keys are separated from reader-facing labels and inverse labels. This prevents a display-name change from changing identity and permits relation vocabulary to participate in the same description, sourcing, audit, and lifecycle machinery as other content.

Layer	Methodological role	Control
Entity	Stable subject identity	Immutable identifier and route key; soft deletion
Entity type	Defines a kind of subject	Data-defined vocabulary; no per-kind physical table
Property definition	Defines an allowed claim	Type-scoped and typed; controlled enumeration where applicable
Fact atom	One assertion about one entity	Status, note, tradition context, source, typed value
Relationship atom	One directed connection	Relation key, status, stance, explanatory note, source
Source entity	Bibliographic and credibility record	Source type, tier, basis, and publication metadata

The design resembles a knowledge graph, but its purpose is not merely graph traversal. It adds editorial state, source accountability, public dispute, and operational provenance to the graph. This aligns with broader knowledge-graph practice while addressing a domain in which perspective and disagreement cannot safely be flattened into a single edge label.[1][2]

4. Evidence Status and Tradition Stance

4.1 Evidence status

Status	Meaning	Required interpretation
Attested	The available sources support the claim	Established for atlas purposes, not metaphysical certainty
Debated	Credible competing scholarly positions exist	The disagreement itself is established
Debunked	The claim has been tested and found false	The rejected claim remains visible as a documented conclusion
Unknown	A documented inquiry reached no defensible answer	Must identify the gap and what evidence could resolve it

Status concerns the evidentiary position of the atlas. It does not describe whether a religious community accepts the claim. 'Unknown' is not a softer synonym for 'debated.' Debate has identifiable positions and evidence; unknown records that inquiry has not produced a defensible answer.

4.2 Tradition stance

Stance	Meaning
Holds	The named body affirms the claim.
Associated	The body is connected to the claim without necessarily asserting it.
Critiques	The body argues against the claim without wholly refusing it.
Rejects	The body refuses the claim.
Held Differently	The body maintains a different account of the same subject; neither account is treated as the default.

The fifth stance is methodologically important. 'Held Differently' prevents non-equivalent accounts from being forced into agreement or opposition. It is appropriate where accounts share a subject but differ in conceptual framing, narrative, causation, or vocabulary. The renderer fails visibly on an unrecognized stance rather than silently treating it as agreement.

5. Sources, Citations, and Reliability Tiers

A source is a first-class entity rather than a free-text footnote. Before a newly introduced source may support claims, the operating rules require the bibliographic fields that honestly exist, a source type, a reliability tier, a basis for that tier, and a reader-facing credibility explanation. Where a field genuinely has no honest value, the absence should be explained rather than concealed.

5.1 Three-tier reliability model

Tier	Operational interpretation	Current sources
Tier 1	Scholarly reference, peer-reviewed work, critical edition, or comparably authoritative primary/custodial record	3,437
Tier 2	Credible specialist, institutional, journalistic, or well-supported secondary source	1,699
Tier 3	Useful but more limited, derivative, partisan, popular, or context-dependent source requiring caution	1,718

The tier classifies provenance and editorial reliability; it does not mechanically decide whether a particular claim is correct. A high-tier source can be misapplied, and a lower-tier source can be appropriate evidence for a community's own contemporary position. The claim-level decision therefore remains distinct from source classification.

As of the August 17 launch-day snapshot, the database holds 6,856 live sources, of which 6,854 carried one of the three tier values. The remaining two carry a value outside that vocabulary, awaiting a verified substitute source; six of the original eight were resolved between this paper's first and second launch-day drafts, either reclassified into a tier or retired in favor of a verified replacement, so this is a shrinking count rather than a static one. The completed basis-note corpus includes recurrent edge cases rather than only ideal examples: state media identified as such, advocacy organizations used for positions they advance, institutions used as primary sources for their own self-description, and syndicated material distinguished from genuinely independent coverage. The launch-day snapshot records 138,802 citations supporting 159,467 public reader-entry fact records. These totals should be read as operational scale, not as proof of scholarly validity.

5.2 Citation completeness

The principal write procedure rejects an attested fact without a source, except for narrowly identified descriptive or structural fields whose sourcing is carried at the entity level. The executable citation gate checks for uncited attested records and permits only exact, documented exceptions. This is stronger than relying on editorial memory, but it does not prove that each citation entails the claim it accompanies. Entailment and interpretive fit still require review.

6. Unknowns, Disagreement, and Open Questions

The atlas treats the boundary of knowledge as publishable information. An unknown-status claim requires an explanatory note recording the inquiry that stalled. A permanent quality gate detects an unknown claim that lacks a linked open question. The open question states why the issue remains unresolved, which disciplines may be needed, and what evidence would settle it.

This changes the editorial incentive. A contributor is not rewarded for manufacturing completion when a source cannot support it. Recording a precise unknown becomes a successful result because it creates a citable research agenda rather than a blank field.

The public open-question population is gated separately from the total internal inventory. At the current snapshot, the system contains 598 open questions, of which 571 are marked public. Resolved and withdrawn questions remain distinguishable in the underlying record.

7. Editorial Production and Review Workflow

The operational workflow was developed through repeated audits and recorded decisions. It is designed for a mixed environment in which humans direct the project and automated agents assist with research, structuring, review, and implementation. The workflow is not equivalent to conventional journal peer review; it is an internally governed production and verification system.

1. Scope the subject and check standing decisions that apply to it.
2. Identify candidate sources and complete their bibliographic and credibility records.
3. Open an auditable work session with a named actor and pass tag.
4. Stage or claim the affected entities so concurrent work cannot overwrite them.
5. Write through guarded procedures that enforce vocabulary, source, status, and value rules.
6. Read the resulting records back on a fresh connection; procedure output alone is not accepted as proof.
7. Perform a second-pass review, including adversarial checks for invented facts, misplaced certainty, internal process language, cultural restrictions, and citation fit.
8. Run the relevant executable gates and verify the rendered public surface.
9. Record unresolved questions, dissent, or an owner decision in the appropriate append-only register.

A key operating principle is class-based correction. A defect discovered on one page is treated as one instance of a potentially systemic pattern until a wider inventory proves otherwise. The intended outcome is not merely to repair the visible instance but to determine whether the same failure mode appears elsewhere and whether a permanent gate can detect its recurrence.

8. Auditability, Quality Gates, and Reproducibility

8.1 Audited changes

Every governed write is associated with an audit session identifying the actor, optional authenticated user, pass tag, and note. Insertions are logged at record level; updates and soft deletions are recorded at column level with old and new values. The audit operation is executed inside the database write procedures so it covers both application users and automated maintenance processes.

As of August 17, launch day, the database contains 27,759 audit sessions and 886,694 audit-log records. These counts show use of the mechanism, not correctness of every recorded decision. Their methodological value is that a later reviewer can reconstruct who changed what, when, and under which work pass.

8.2 Decision register

The project maintains an append-only owner-ruling register. Decisions carry a summary, scope tags, source context, status, and an explicit supersession link when replaced. A new directive that conflicts with a standing decision must identify the conflict and receive explicit confirmation rather than silently overwriting institutional memory. The current register contains 654 rulings, including 517 marked standing.

8.3 Executable quality gates

Rules that can be tested reliably are registered as executable gates with a failure predicate, allowed count, and documented reason for any exact exception. The registry mixes release gates with censuses and worklists; counts in this paper therefore use the explicit gate flag, not the number of registry rows. On that basis, the August 17 snapshot contains 45 executable release gates among 71 registry records. Examples detect uncited established claims, facts with invalid value cardinality, unknowns without open questions, unresolved database references, relation-vocabulary defects, content-schema drift, internal process language leaking into public prose, and critical stances with neither explanation nor source.

LIMIT OF AUTOMATION A green gate run establishes that its registered predicates passed. It does not establish that every relevant predicate has been conceived, registered, or correctly defined. The playbook records several occasions when a gate list itself was incomplete; the registry is therefore treated as a governed object subject to review.

8.4 Reproducible workflows

New user-facing workflows are required to have an inventory describing actors, seed data, steps, observable outcomes, side effects, and cleanup. Where a simulator exists, it seeds unmistakable test

records, exercises the workflow end to end, verifies observable results, and cleans the records up provably. This provides operational reproducibility, although it should not be confused with replication of scholarly findings.

9. Cultural Restrictions and Rights-Sensitive Material

The atlas distinguishes scholarly description of a tradition from reproduction of restricted practice. Restrictions are represented as cited properties within the same content model rather than as an informal blacklist. A restriction can affect report-generation channels without removing a public scholarly description of the subject.

The governing direction is fail closed: absence of a clearance record is not treated as clearance, and only an explicitly publishable rights determination permits restricted material to pass a channel. Rights determinations identify the affected entity or type, channel, property where relevant, rights holder, original work, source, and a reader-facing placeholder when material is withheld.

The methodology is informed by the broader CARE Principles for Indigenous Data Governance, particularly the insistence that technically open or reusable data can still implicate collective authority, responsibility, ethics, and benefit.[3] Divinity Atlas does not claim that its current implementation fully satisfies CARE or that one restriction model can represent every living community. External review by affected communities remains necessary.

10. Public Challenge, Cases, Corrections, and Dissent

10.1 Challenge and appeal

A registered reader may challenge a specific claim. Escalation requires a specific, checkable source and a description of what should be checked. A qualifying appeal identifies the exact fact or relationship under review rather than merely naming a page or topic.

10.2 Public case file

A case advances through Awaiting Evidence, In Review, Decided, and Closed. Each examined item receives its own case-file record naming what was reviewed, what it showed, its source, and the role of the filer. A decision requires both a concise judgment summary and full reasoning. The disposition vocabulary is upheld, revised, overturned, or declined.

10.3 Dissent and precedent

A reviewer, independent internal lane, or challenger may file a cited dissent. A challenger is publicly named; internal reviewers are identified by role. Cases may cite earlier cases as precedent. These

structures preserve disagreement without forcing the final public page to pretend that adjudication erased every objection.

10.4 Corrections

Material changes that reverse what a page asserted are intended to produce a public correction narrative explaining the previous claim, the new evidence, and why the record changed. The underlying audit history remains intact. Public correction counts begin at the configured go-live date, so the count at launch is correctly zero even though pre-launch correction work exists in the private audit history.

At the snapshot date, the system holds eight cases, eleven case-evidence items, and five dissents. Six cases are closed and two are decided; six dispositions upheld the record and two revised it. This is an operating record, not yet a statistically meaningful validation sample.

10.5 Launch-gated claim identifiers

Stable citation identifiers for individual claims and citable page versions are implemented under owner rulings 378 through 380. They are intentionally gated behind the public go-live date rather than exposed as pre-launch history. This establishes a defined beginning for the public citation record: launch-time state becomes the first publicly citable version, while subsequent changes can be located against stable identifiers and dated versions. The launch gate is therefore part of the methodology, not an unfinished identifier design.

11. Automation, Human Judgment, and AI Disclosure

Automated agents have participated extensively in research assistance, data transformation, code production, audits, and second-pass reviews. Large language models also assist with portions of personalized interpretive report prose. Deterministic calculations, including astronomical positions and system-specific arithmetic, are implemented separately in code and are not delegated to prose generation.

The system's safeguards reduce several known automation risks: invented citations are prohibited; writes require real source identifiers; controlled vocabularies fail loudly; every change is auditable; read-back is required; class-wide audits follow discovered defects; and public challenge remains available. These controls do not remove model bias, source-selection bias, hallucination risk during research, or the possibility that multiple agents share the same training-data assumptions.

DISCLOSURE STANDARD The defensible claim is that automation operates inside a governed, auditable editorial system. It is not defensible to present internal agent review as external scholarly peer review or to imply that database constraints validate the truth of cited content.

Human authority presently includes the project owner, administrative reviewers, and the policies recorded in the rulings register. The proposed next phase is to add identifiable external scholars and community advisers with declared expertise, conflicts, review scope, and authority to recommend revision without requiring endorsement of the project as a whole.

12. Current Scale and Operational Evidence

The following figures are a live database snapshot from August 17, 2026, launch day. They describe scale and process coverage, not academic validation.

Corpus counts in this table use the public reader-entry definition rather than raw ontology-table row counts. 'Entries' excludes source records, articles, open questions, and vocabulary or other supporting entities that share the entity table. 'Public fact records' and 'documented public relationships' count live claims and edges attached to that public reader corpus, excluding internal, deleted, supporting, and non-reader records. 'Citations' counts citation links supporting those public reader-entry facts and relationships rather than every self-citation or internal provenance link stored by the system. Raw table counts are therefore expected to be substantially larger. Publication-day figures should be re-extracted in the same pass used for the launch press release so both documents report one dated snapshot.

Measure	Snapshot
Entries	16,156
Entry kinds	237
Public fact records	159,467
Citations	138,802
Sources	6,856
Documented public relationships	69,965
Articles	29,914
Religions	181
Disciplines represented	116
Open questions	598 total open; 571 public
Owner rulings	654 total; 517 standing
Quality registry	71 total records; 45 executable gates
Audit trail	27,759 sessions; 886,694 audit records

Measure	Snapshot
Public cases	8

13. Limitations and Validation Agenda

13.1 Present limitations

- The corpus has not undergone representative external scholarly review across all 237 entry kinds.
- Internal second-pass review can separate tasks and roles but does not establish institutional independence or subject-matter credentials.
- A citation requirement tests presence, not entailment, correct interpretation, or adequacy of source diversity.
- Source tiers are operational editorial judgments and require published calibration examples and inter-rater study.
- Coverage volume is uneven across traditions, regions, languages, periods, and source ecosystems.
- English-language and digitally accessible sources are likely overrepresented.
- Public challenge participation can reflect access, confidence, and community resources rather than the actual distribution of errors.
- The correction and case sample is too small and too new for outcome claims.
- Rights and cultural-restriction decisions require continuing consultation with affected communities.
- The model's categories may reproduce assumptions embedded in existing scholarship and in the project's own architecture.

13.2 Proposed validation program

10. Recruit an external advisory group spanning religious studies, theology, anthropology, history, library science, digital humanities, knowledge representation, and community-based expertise.
11. Draw a stratified sample across entry kinds, traditions, source tiers, statuses, and geographic regions.
12. Have at least two qualified reviewers independently assess claim accuracy, citation entailment, status assignment, stance assignment, and cultural framing.
13. Measure agreement and report both raw agreement and an appropriate chance-corrected statistic, while preserving qualitative disagreement.
14. Publish the sampling frame, review instrument, reviewer qualifications, conflicts, adjudication procedure, and all non-sensitive results.
15. Evaluate the launch-gated stable claim identifiers and citable page versions, and provide a downloadable research snapshot.
16. Repeat the study after remediation and publish changes rather than replacing the first result.

This program would allow the project to move from a well-governed internal methodology to externally examined evidence about accuracy, consistency, and bias. It would also make the white paper itself falsifiable: claims about the method could be compared with measured outcomes.

13.3 Pre-external internal validation study

A useful validation study does not need to wait for an external advisory group. Before external recruitment, the project can draw a stratified sample of several hundred claim atoms and assign them to independent re-verification lanes that did not author the sampled records. The study should measure reviewer agreement, citation presence and entailment, status and stance agreement, correction frequency, and disposition outcomes preserved in the audit trail. Results should be reported by source tier, entry kind, tradition, geographic region, and risk category rather than as a single undifferentiated accuracy rate.

The existing audit record supplies candid test material. Recent operations recorded duplicate incidents that were corrected and followed by stronger prevention procedures; proposed charter premises that were refuted when lanes checked the database before acting; and conflict records whose missing knowledge was represented as formal unknowns rather than silently inferred. These examples should be sampled as failure-recovery cases, not promoted as evidence of accuracy. Their value is that they expose whether the control system detects, records, and learns from error.

Coverage bias can be examined using the live axis-depth mechanism, which measures where the corpus is thin across defined dimensions. A published validation report should use those measurements to describe undercoverage, guide stratified sampling, and avoid allowing heavily developed areas to conceal weakly represented traditions, regions, periods, or claim types.

14. Conclusion

Divinity Atlas is built on the premise that uncertainty and disagreement should be modeled, not edited away. Its architecture separates evidence status from tradition stance, represents sources and open questions as first-class records, preserves operational provenance, and makes contested claims addressable through a public case process. Its quality system combines guarded writes, audit history, executable gates, adversarial review, and workflow simulation.

Those mechanisms are meaningful, but they do not confer scholarly authority by themselves. The project should be judged by whether its citations support its claims, whether its categories fairly represent the communities described, whether its uncertainty labels are applied consistently, and whether external review produces visible correction. This paper therefore presents the methodology as an inspectable framework and an invitation to test it, not as proof that the atlas has already earned academic consensus.

Appendix A. Technical Control Map

Risk	Primary control	Residual limitation
Invented or missing citation	Source identifier required by guarded write procedure; uncited-attested gate	A present citation may still be misapplied
Silent overwrite	Audited procedures, staging claims, fresh-connection read-back	Authorized but mistaken updates remain possible
Vocabulary drift	Controlled enumerations, database checks, visible failure on unknown stance	The controlled vocabulary itself may be incomplete
Unknown presented as fact	Status vocabulary, explanatory note, open-question linkage gate	Editorial judgment still determines when inquiry is sufficient
Disagreement flattened	Separate stance model including Held Differently	Complex positions can exceed a five-value vocabulary
Internal process prose exposed	Pattern-based internal-language gate and rendered-page review	Novel leak forms may evade known patterns
Restricted material reproduced	Cited restriction records and fail-closed channel rights	Community expectations may differ from recorded policy
Decision memory loss	Append-only rulings with supersession links	Owner-centered governance remains a concentration of authority
Correction concealed	Audit history, correction narrative, public case docket	Some prose edits without status transition require explicit treatment
Workflow regression	Scenario inventory and simulator with seeded cleanup	Simulation coverage depends on inventory completeness

Appendix B. Proposed External Scholarly-Review Protocol

B.1 Review unit

One review unit consists of an entry, a selected claim atom or relationship atom, its visible context, all displayed citations, status and stance labels, explanatory notes, and the version timestamp. Reviewers should assess the atom in context rather than judging a detached database row.

B.2 Reviewer instrument

Dimension	Reviewer question	Response
-----------	-------------------	----------

Dimension	Reviewer question	Response
Identity	Is the subject correctly identified and disambiguated?	Yes / partial / no / cannot assess
Accuracy	Does the claim fairly state what the evidence supports?	Yes / minor revision / major revision / unsupported
Entailment	Does each citation support the specific claim made?	Direct / indirect / weak / none
Status	Is attested, debated, debunked, or unknown appropriate?	Agree / alternate status / cannot assess
Stance	Is the tradition-specific stance represented without imposing a baseline?	Agree / revise / not applicable
Framing	Is the account fair, proportionate, and free of avoidable outsider bias?	Accept / concern / serious concern
Restriction	Does the treatment respect applicable community and rights constraints?	Accept / consult / withhold

B.3 Sampling

Use stratified random sampling with deliberate oversampling of debated, debunked, unknown, restricted, living-tradition, lower-tier-source, and cross-tradition records. Publish both weighted corpus estimates and unweighted high-risk results. Do not select only showcase entries.

B.4 Governance

Record reviewer expertise, institutional affiliation where applicable, conflicts, community relationship, and review scope. Permit reviewers to publish dissent from the final disposition. Escalate living-community and restricted-material concerns to appropriate community expertise rather than resolving them solely through general academic credentials.

B.5 Reporting

Report the sample, exclusions, missing data, disagreement, correction rate, citation-entailment results, and status-assignment agreement. Publish a versioned remediation ledger linking each accepted finding to the resulting correction, open question, case, or documented decision not to change.

References

- [1] Hogan, A., Blomqvist, E., Cochez, M., et al. (2021). Knowledge Graphs. *ACM Computing Surveys*, 54(4), Article 71. <https://doi.org/10.1145/3447772>
- [2] Vrandečić, D., and Krotzsch, M. (2014). Wikidata: A Free Collaborative Knowledgebase. *Communications of the ACM*, 57(10), 78-85. <https://doi.org/10.1145/2629489>
- [3] Carroll, S. R., Garba, I., Figueroa-Rodriguez, O. L., et al. (2020). The CARE Principles for Indigenous Data Governance. *Data Science Journal*, 19, 43. <https://doi.org/10.5334/dsj-2020-043>
- [4] Wilkinson, M. D., Dumontier, M., Aalbersberg, I. J., et al. (2016). The FAIR Guiding Principles for scientific data management and stewardship. *Scientific Data*, 3, 160018. <https://doi.org/10.1038/sdata.2016.18>
- [5] World Wide Web Consortium. (2013). PROV-O: The PROV Ontology. W3C Recommendation. <https://www.w3.org/TR/prov-o/>
- [6] International Council of Museums, CIDOC CRM Special Interest Group. CIDOC Conceptual Reference Model. <https://www.cidoc-crm.org/>
- [7] Divinity Atlas. Active operational playbook, owner-ruling register, quality-gate registry, audit schema, case system, and workflow inventory. Internal project sources inspected August 17, 2026.

Review Note

This draft is intentionally conservative. It distinguishes implemented controls from proposed external validation, avoids treating internal agent review as academic peer review, and reports database counts as a dated snapshot. Reviewers are invited to challenge the accuracy of the process description, identify omitted controls or limitations, and recommend changes before public release.