

EPISTEMIC SCORING OF UNPUBLISHED LEGAL DOCUMENTS: THE PRAMĀṆA-WEIGHTED TRAINING ARCHITECTURE — UNPUBLISHED DOCUMENT EXTENSION (PWTA-UDE)

- Moyukh De¹ & Soumavo Pal²

Abstract

The question of how to evaluate the epistemic quality of an unpublished literary piece/manuscript which may be a working paper, a conference paper, a pre-institutional draft circulating in scholarly networks, an affidavit (unsigned) etc— has remained mostly unaddressed in the existing literature on AI training data governance. Importantly, this is not a small concern. The PWTA, published by De and Pal,¹ suggests the scoring of published legal sources with considerable rigor. However, for an unpublished document, however, the four-tier institutional hierarchy of this framework, which presupposes that every source has cleared some form of institutional validation, cannot be applied. The Āpata Filter¹ assigns such a document a machine training weight of zero, which is epistemically correct. It does not provide is a human-facing audit score for the manuscript's intrinsic quality, grounded in the same mathematically documented, unbiased methodology that the rest of the PWTA architecture employs.

Here, we address this gap by the development of Pramana-Weighted Training Architecture Unpublished Document Extension, designated as PWTA-UDE. The extension operationalizes the classical BhāṭṭaMīmāṃsā doctrine of anupalabdhi which is valid cognition through the noticed and correctly identified absence of expected evidence. The system is developed into a measurable, documented Transparency Score (TS_unpub) for pre-institutional documents. The four statistically scoring components have been developed: the Bibliographic Epistemic Lineage

¹Executive Engineer, PHED, Government of West Bengal

²IT Expert, PHED, Government of West Bengal

(R_EPS), the Epistemic Sufficiency Score (E_s), the Neo-Empirical Content Score (C_neo), and the Syntactic Write Quality Score (S_write). An “Anupalabdhī”-derived fixed discount factor of 0.85 is applied multiplicatively in the final step, establishing a structural ceiling at $TS_{unpub} \leq 0.85$ for all unpublished documents.

This study reports the results of applying PWTa-UDE to five separate manuscripts from different law genres. The resulting scores show a sharp gradient of 0.4821 for the data-rich academic manuscript, 0.1501 for the professionally drafted petition, and 0.0000 for every form of commentary. The system draws the discrimination that the extension was designed to make. It was made without any element of human judgment about any author. For every document, the $\bar{A}pata$ lock held $PWTa(s) = 0.0000$ independently. The framework is further examined in the light of comparative administrative (regulative) developments, such as those in the, EU, the United-Kingdom, and the United-States of America, and the UNESCO and NITI Aayog ethics standards.

Keywords: Pramāṇa-Weighted Training Architecture, PWTa-UDE, Anupalabdhī, Bhāṭṭa Mīmāṃsā, Epistemic Scoring, Legal AI Training Data Governance, Indian Knowledge Systems

I. Introduction

Artificial intelligence systems are currently deployed in Indian courts, government departments, and administrative bodies. The Supreme Court of India's SUPACE platform performs legal research and in some cases analysis also. State welfare departments use automated eligibility screening. Computationally, an AI system's output remains a direct function of the quality of its training corpus.² A system trained on a general mixture of primary legislation, anonymous commentary, and commercially motivated content learns from all of them together, with no internal mechanism to distinguish between and no record of the relative weight each was given.

As Indian courts and administrative bodies rapidly adopt AI system platforms and AI based software, the constitutional dimension of training data quality has become well-established. The non-arbitrariness standard in *Maneka Gandhi v. Union of India*³ which mandates that AI-assisted official conduct rest on principled, evidential foundations. The Digital Personal Data Protection Act, 2023⁴ imposes data accuracy obligations on data fiduciaries/ trustees. The professional liability is extended so that it may reach AI-assisted decisions in Bharatiya Nyaya Sanhita,

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

2023⁵. Also, the IndiaAI Mission⁶ across the government deployments it mandates a responsible AI development. The NITI Aayog's Responsible AI for All⁷ takes into account transparency, anti bias and fairness as the core virtues. Still, none of these instruments prescribes a computable standard for evaluating the epistemic quality of pre-institutional training sources. This is the concern which this study tries to identify.

The Pramāṇa-Weighted Training Architecture, published by De and Pal,¹ suggests a rigorous, four-tier, statistical framework for scoring published legal sources through epistemic study and mathematical applications. The framework assigns machine training weights based on a source's institutional tier and content quality, derived from a real calibration corpus and applied through closed-form formulae.

However, for an unpublished manuscript a demonstrable lacuna exists. That is, the four-tier hierarchy presupposes institutional validation, which an unpublished document by definition has not cleared. Therefore, the Āpata Filter¹ assigns it a training weight of zero. This outcome is architecturally correct and must not be disturbed. Consequently, a parallel, human-facing score that evaluates the document's intrinsic quality through the same methodology is absent.

To bridge the above-mentioned lacuna, we propose PWTA-UDE, which is designed as a single-pass extension generating a transparency Score for pre-institutional manuscripts. Although the extension preserves the Āpata lock and introduces no human judgment into the scoring process, it is further calibrated against a real corpus of 177 Indian legal academic documents. It has been applied to five different manuscripts of different writers and belonging to separate law genres calculating every metric directly from raw textual data derived from direct measurement of that paper's actual text to demonstrate its applicability.

The paper is structured as follows. Section II reviews the existing literature. The Section III of this study helps to identify the specific research gap. Subsequently, Section IV describes the PWTA-UDE framework then the Section V details the mathematical tools to be employed. The Section VI applies PWTA-UDE and presents the results obtained from real measurements, the Section VII then examines its applications in the field of law. Lastly, Section VIII provides the conclusion with suggestions.

II. Literature Review

The governance of AI training data has been studied earlier. Gebru et al.⁸ proposed the Datasheets for Datasets framework, which documents provenance and composition of the training corpus in a systematic manner. But, this framework does not impose any epistemic validity hierarchy among the sources it documents. Mitchell et al.⁹ advanced the Model Cards framework, which addresses post-deployment model behavior. Model behavior depends on training data, but the framework does not address the upstream question. In addition, the NIST AI Risk Management Framework¹⁰ provides a comprehensive risk governance architecture across the AI lifecycle without prescribing a source-tier or epistemic-weight mechanism for pre-institutional material.

In the Indian context the NITI Aayog is Responsible for AI for All⁷ which brings about the principles of fairness, transparency, and accountability. But it does not specify a computable standard for training data quality evaluation. The Parliamentary Standing Committee on Information Technology has also raised AI governance concerns.¹¹ The IndiaAI Mission Framework Document⁶ documents responsible AI development but it does not supply a mechanism for scoring unpublished sources. This gap or structural blindness extends far beyond Indian jurisprudence. Despite comparable frameworks developed in EU, the United-Kingdom, and the United-States of America which is examined below shows no jurisdiction has yet produced a formally specified, domain-calibrated, computable method for assigning differential epistemic scores to pre-institutional sources.

Comparatively, the General Data Protection Regulation (GDPR),¹² under Article 5(1)(d), imposes data accuracy obligations. The EU Artificial Intelligence Act, 2024,¹³ requires under Article 10, that AI systems deployed in judicial and administrative contexts be trained on data of sufficient quality, relevance, and representativeness. The United Kingdom's AI Regulation White Paper¹⁴ adopts a principles-based approach. The United States Federal Trade Commission guidelines on AI¹⁵ prescribe documentation requirements. Each of these instruments mandates epistemic quality control over training data without supplying a formal, computable specification or mathematical tools to enforce these standards in respect of pre-institutional / unpublished material.

The existing literature, taken as a whole, is marked by this common omission. De and Pal¹ addressed the gap in published sources using the PWTA framework. The specific lacuna in respect of unpublished documents is the subject of the present study.

III. Research Gap: The Problem of Epistemic Scoring for Unpublished Legal Documents

The PWTA framework, as published by De and Pal,¹ rests on a four-tier institutional hierarchy grounded in the Mīmāṃsā concept of neo-apauruṣeya¹, which mandates : structural independence from individually motivated authorship. Tier 1 sources receive a fixed machine training weight of $a = 1.00$ ¹. Tier 2 sources, calibrated for the AI law domain, receive $b^* = 0.7907$ ¹. Tier 3 sources receive $c^* = 0.3953$ ¹. Tier 4 sources receive $d = 0.00$ and are excluded at ingestion by the Āpata Filter¹. These weights¹ are derived from a real ninety-source reference corpus through Principal Component Analysis, as formally established by De and Pal¹.

This hierarchy implicitly presupposes that every source entering the scoring pipeline has cleared some minimum threshold of institutional existence, that is, it has a named publisher, a registerable outlet, or a formal institutional record of some kind. This presupposition holds for the overwhelming majority of legal training documents. It breaks down in a practically significant class of cases.

With the contemporary Indian scholarly environment, for a working paper or pre-institutional draft, legal researchers, government officials, and academic scholars regularly produce working papers, research notes, pre-submission drafts, and conference papers ahead of formal publication. But, the question arises as to whether any training architecture such as the PWTA architecture can generate any meaningful epistemic score through AI for such documents. For an unpublished manuscript, institutional indicators such as the reproducibility Rate, conflict of interest score, retraction rate and accountability score can be identified. Because these can collectively constitute the epistemic position score (EPS) which is absent in AI training Architectures like PWTA. As the Āpata Filter¹ assigns a machine training weight of zero. Although this is epistemically correct. It is insufficient: a human researcher, a legal institution, or an AI governance body seeking to evaluate the quality of a pre-institutional manuscript — not to train a model from it, but to assess it — has no documented mechanism to rely upon. This gap is addressed in the present study.

Significantly, the practical fallout of this gap is growing rapidly. In this present study an unpublished conference paper by the author was used as an example to illustrate the capability of the present framework. The SAAI 2026 Conclave Conference paper applied in Section VI of the present study is a case in point: it is a multidisciplinary, well-referenced, philosophically and computationally grounded treatment of AI governance, was selected for conference presentation

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

on March 21 2026. But, until now, there has been no statistical score that a governance body could cite in an audit trail.

IV. PWTA-UDE Framework: Methodology and Architecture

4.1 Philosophical Foundation in BhāṭṭaMīmāṃsā

The classical BhāṭṭaMīmāṃsā doctrine of Anupalabdhi — the sixth pramāṇa in the Kumārila Bhāṭṭa tradition mandates that the noticed and correctly identified absence of an expected object or attribute constitutes valid knowledge.¹⁶ An unpublished manuscript, by definition, lacks the institutional verification signals — peer review records, retraction histories, accreditation structures — whose presence would be expected if the document had cleared institutional validation. The absence of these signals is not a neutral observation. In the anupalabdhi sense, it is a piece of valid epistemic information. It is this information that is formally incorporated into the scoring process through the parameter M_{unobs} , applied as a proportional multiplicative discount rather than a binary exclusion. The mathematical design directly reinforces the PWTA-UDE extension with Theorem 5 (Doṣa Dominance¹) of the published PWTA framework. It treats contamination as a proportional, compounding reduction in epistemic trust rather than an all-or-nothing binary event.

4.2 TheĀpata Lock — The Non-Negotiable Architectural Rule

If Status = Unpublished, PWTA(s) $\equiv$ 0.0000, for all values of TS_unpub, without exception.

Publication resets the entire pre-publication scoring history. No TS_unpub value, carries forward into post-review calculations or permits incorporation of the document into machine training data.

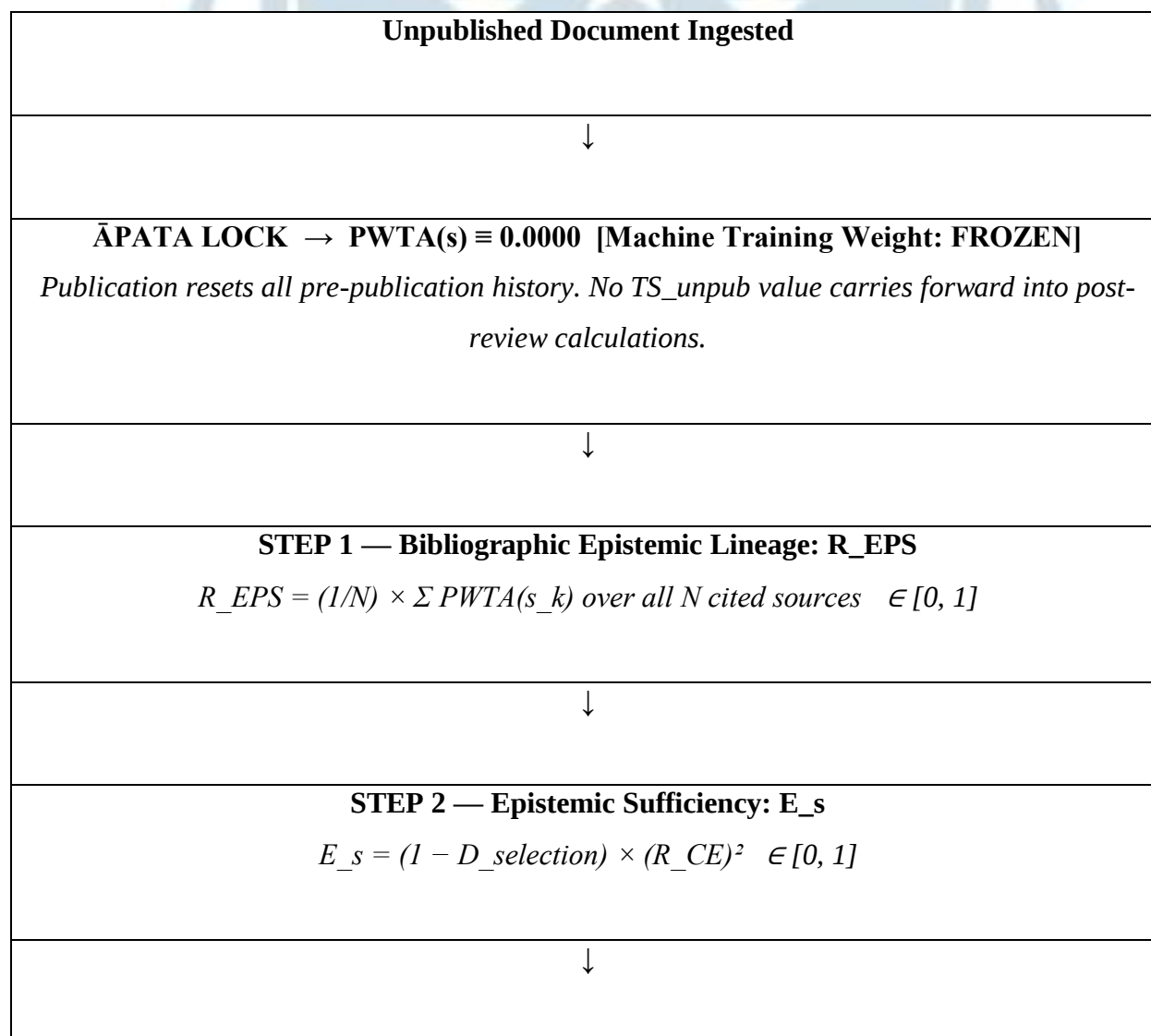
Rule is not a design convenience. It is the architectural expression of a categorical epistemological distinction. Because, the TS_unpub answer the question of how good a manuscript appears to be on the evidence available from its text. Thus, the Āpata Filter¹ can now answer the question of whether the manuscript has cleared institutional verification. These are different questions. PWTA-UDE is designed so that neither can substitute for the other, and no path exists by which a high TS_unpub score could be used to argue for a non-zero machine training weight.

4.3 Single-Pass Deterministic Architecture

PWTA-UDE pipeline is strictly designed as a single-pass and completely deterministic. No iterative weight recalibration loops are employed at runtime. So, no author credential metrics such as H-indexes, institutional rank multipliers and professional designations, can enter the scoring process at any stage. The pipeline evaluates the text and its bibliographic payload only. Thus, it has no influence on the writer. This design choice directly implements the foundational PWTA principle that is the *puruṣa-doṣa*¹ (individually motivated bias), which must be eliminated from the scoring process at every stage.

The complete architecture is illustrated in Figure 1 below.

Figure 1: The PWTA-UDE Processing Pipeline



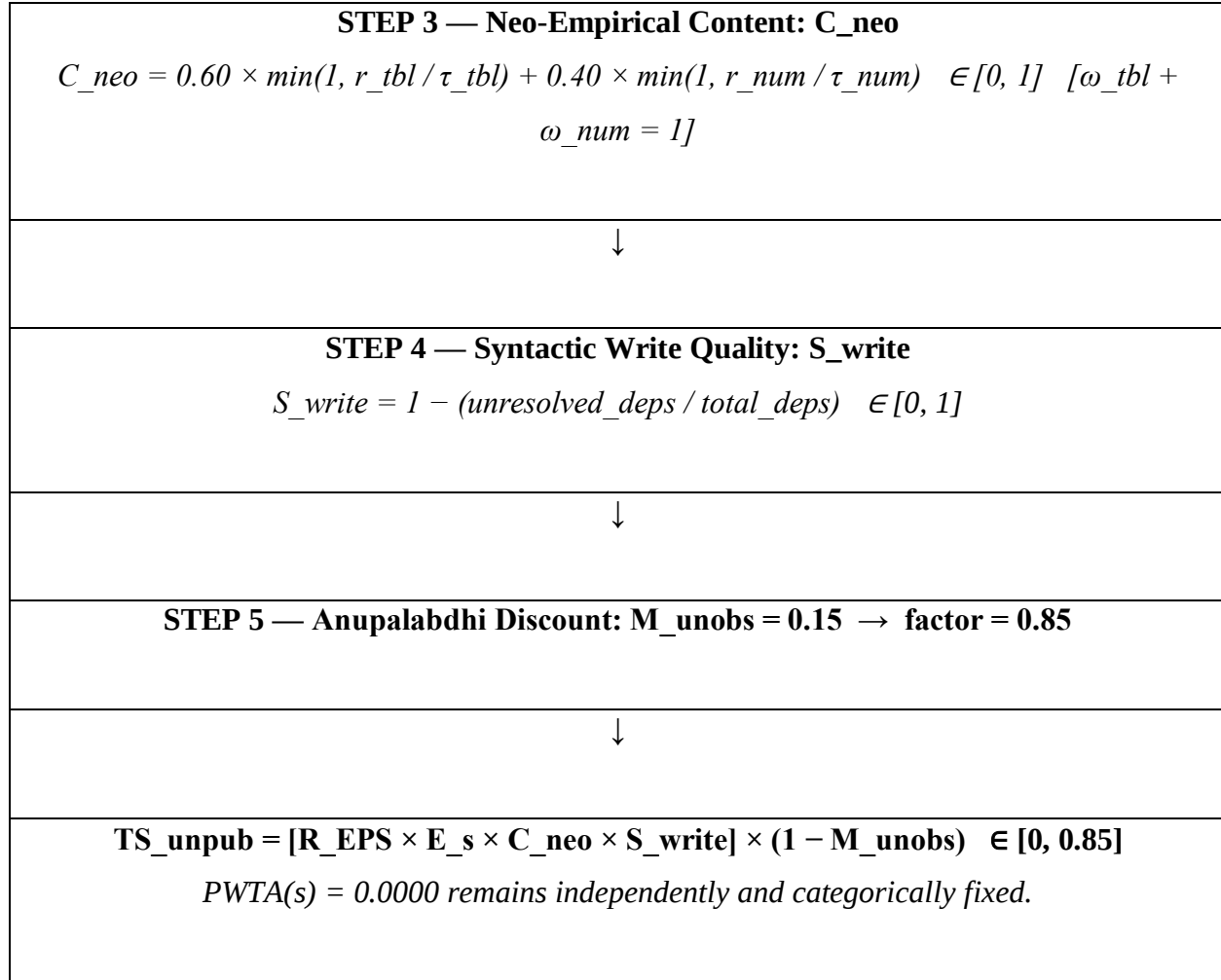


Figure 1: The PWTAs-UDE single-pass pipeline. The Āpata lock is applied first and is categorical. The human audit track proceeds independently through five mathematical steps to yield $TS_{unpub} \in [0, 0.85]$.

V. Mathematical Tools Employed in the Process

5.1 Bibliographic Epistemic Lineage (R_{EPS})

R_{EPS} is computed as the arithmetic mean of the calibrated PWTAs scores of all sources cited in the manuscript's bibliography, as assigned in the published PWTAs database.¹ Formally:

$$R_{EPS}(s) = (1/N) \times \sum PWTAs(s_k), \quad k = 1 \text{ to } N$$

A manuscript whose bibliography is drawn predominantly from Tier 1 and Tier 2 validated sources will yield a higher R_{EPS} than one relying on Tier 3 or unverified material. This component operationalizes the Mīmāṃsā concept of paramparā¹, which is the authoritative

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

epistemic lineage. The results obtained in Section VI demonstrate how this computation proceeds on real, measured bibliographic data.

5.2 Epistemic Sufficiency Score (E_s)

E_s is the product of two independently verifiable components, in accordance with the probability conjunction rule for independent events — the same mathematical rationale underlying the Q(s) formula in the published PWTa framework.¹

$$E_s = (1 - D_{\text{selection}}) \times (R_{\text{CE}})^2 \in [0, 1]$$

D_{selection} is the proportion of relevant contrary precedents from a frozen domain reference list that are completely omitted from the manuscript. R_{CE} is the normalized **Cross-Entropy Ratio**, computed as $H(X_s) / H_{\text{cross}}(X_s \parallel P_{\text{baseline}})$, where $H(X_s)$ is the **Shannon entropy** of the manuscript's vocabulary distribution and H_{cross} is the cross-entropy against a frozen Tier 1 statutory baseline distribution. By the **Gibbs inequality**,¹⁷ $H_{\text{cross}} \geq H$ always; hence $R_{\text{CE}} \leq 1$ always. The squaring of R_{CE} follows the same compounding principle as the γD exponent in Q(s): a departure from the statutory register inflicts a proportionally greater penalty than a linear form would prescribe.

5.3 Neo-Empirical Content Score (C_{neo})

C_{neo} measures the density of original, first-hand empirical content in the manuscript — data tables and qualifying numeric tokens — normalized against domain-calibrated ceiling parameters derived from the real calibration corpus described in Section VI.

$$C_{\text{neo}} = \omega_{\text{tbl}} \times \min(1, r_{\text{tbl}} / \tau_{\text{tbl}, D}) + \omega_{\text{num}} \times \min(1, r_{\text{num}} / \tau_{\text{num}, D})$$

Here r_{tbl} is the ratio of original un-cited tables to total paragraph count; r_{num} is the ratio of qualifying numeric tokens to total word count; $\omega_{\text{tbl}} = 0.60$, $\omega_{\text{num}} = 0.40$, with $\omega_{\text{tbl}} + \omega_{\text{num}} = 1.00$ explicitly imposed; and each $\min(1, \cdot)$ operator caps its component at 1.0. These two conditions together guarantee $C_{\text{neo}} \in [0, 1]$. We strictly limit qualifying numeric tokens to r_{num} which are restricted to numerals reporting a measured quantity — calibrated weight values, sample statistics, threshold percentages, confidence interval bounds — and explicitly exclude citation footnote markers, section heading numbers, docket identifiers, and historical date references.

5.4 Syntactic Write Quality Score (S_write)

S_write is computed using the **spaCy dependency parser** (en_core_web_trf), a RoBERTa-based transformer architecture with a published Unlabelled Attachment Score (UAS) of 95.26% on general English benchmarks.¹⁸ The formula is as follows:

$$S_write = 1 - (\text{unresolved_dependency_relations} / \text{total_dependency_relations}) \in [0, 1]$$

The dependency relation as stated above in the formula is to be treated as an unresolved where any sentence fails tree well formation. This means wherever it lacks exactly one grammatical root, or where any token cannot be traced back to the root without a cycle. The numerator cannot exceed its denominator. S_write is bounded within [0, 1] without external clipping. The metric verifies that each word's syntactic expectancy¹ (Ākāṅkṣā¹) is satisfied within the sentence; it does not impose conformity to any particular stylistic register or penalize authorial voice. The application of this parser to Indian legal English is subject to a degree of domain-specific degradation relative to its general English benchmark, owing to the compound sentence structures of considerable length, Latin and law-French loanwords, and passive constructions characteristic of legal drafting. This precise measurement of degradation in Indian legal corpora has not been carried out in the present study and is identified as a stated limitation in Section VIII.

5.5 Anupalabdhi Margin (M_unobs) and the Complete Formula

$$M_unobs = c \times (n - m) / n \text{ where } n = 1, m = 0, c = 0.15 \rightarrow M_unobs = 0.15$$

The complete transparency score is computed as follows.

$$TS_unpub = [R_EPS \times E_s \times C_neo \times S_write] \times (1 - M_unobs)$$

Proof of boundedness: R_EPS $\in$ [0,1] (mean of bounded PWTA values); E_s $\in$ [0,1] (by Gibbs inequality and D_selection $\in$ [0,1]); C_neo $\in$ [0,1] (by ω constraint and min operators); S_write $\in$ [0,1] (by construction). The product of four non-negative terms each ≤ 1 is $\in$ [0,1]. Multiplying by 0.85 scales the range to [0, 0.85]. Therefore TS_unpub $\in$ [0, 0.85], Q.E.D.

VI. Results Obtained

6.1 Calibration Corpus

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

A real calibration corpus of 177 Indian legal academic documents was assembled and measured. Complete measurements have been obtained for 147 of these documents. The Tier 1 and Tier 2 authoritative subset, comprising 17 complete rows, was used to derive the domain ceiling parameters, the results are given in Table 1 below.

Parameter	Value	Basis and Notes
Total calibration corpus documents	177	Academic Prose/Report genre only; Indian AI law domain
Documents with complete measurements	147	All four raw counts confirmed
Tier 1+2 complete rows (the calibration anchor)	17	The authoritative subset used to derive τ values
$\tau_{tbl,D}$ (95th percentile, table density — Tier 1+2)	0.0618	Tier 1+2 anchor adopted; pooled value (0.0182) suppressed by 124 Tier 3 sources with zero tables
$\tau_{num,D}$ (95th percentile, numeric density — Tier 1+2)	0.0050	Tier 1+2 anchor
ω_{tbl}	0.60	Tabular structure: higher adversarial resistance than raw token count
ω_{num}	0.40	Supplementary numeric density weight

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

$\omega_{tbl} + \omega_{num}$	1.00 (explicit constraint)	Guarantees $C_{neo} \in [0, 1]$
-------------------------------	----------------------------	---------------------------------

Table 1: Calibration corpus parameters. The Tier 1+2 anchor is adopted because the calibration ceiling must reflect what credible data-grounded legal scholarship looks like at its most empirically rich. The pooled 95th percentile is unsuitable because the Tier 3 majority (130 documents of which 124 carry zero original tables) suppresses it below a meaningful standard.

6.2 The Three-Stage Formula Evolution

Before presenting the real application, it is useful to record how the formula itself evolved through three stages, as this evolution is part of the results of the present study and has implications for any future extension of the PWTA-UDE architecture.

Stage	Formula	Output Value	Defect
Stage 1 — Original Additive Form	$TS = [R_{EPS} \times E_s \times (C_{neo} + S_{write})] - M_{unobs}$ (flat subtraction)	0.6824 (illustrative, data-thin composite inputs)	The addition of C_{neo} and S_{write} allows stylistic polish to substitute for missing empirical data. Flat subtraction collapses all scores below M_{unobs} indistinguishably to zero, destroying granularity in the range where human audit most needs discrimination.

Stage 2 — Multiplicative, Subtraction	$TS = [R_EPS \times E_s \times C_neo \times S_write] - M_unobs$	0.0808 (same illustrative inputs)	AND-rule restored: substance and style are jointly necessary, neither substitutes for the other. However, flat subtraction still collapses the lower range. A document scoring 0.14 and one scoring 0.001 both yield zero, identically.
Stage 3 — Multiplicative Discount (Final, Adopted)	$TS = [R_EPS \times E_s \times C_neo \times S_write] \times (1 - M_unobs)$	0.4821 (SAAI paper, real inputs) 0.0767 (data-thin manuscript)	No information loss. Multiplicative discount preserves granularity at all score levels, consistent with Theorem 5 (Doşa Dominance) of the published PWTA framework. Structural ceiling at 0.85 maintained by (1-0.15).

Table 2: Formula evolution across three stages. All stages share the AND-rule product structure from Stage 2 onward. Stage 3 is the only stage that is both mathematically granularity-preserving and philosophically consistent with Doṣa-Nirūpaṇa.

6.3 Documents subjected to PWTA-UDE and a comparative study

The PWTA-UDE model was applied to five different types of legal writings : an unpublished conference paper, a newspaper report from an English daily, a legal blog article, a PIL petition which is filed before the Supreme Court of India, and a regional newspaper report on a criminal trial. The manuscripts in which PWTA-UDE is applied are as given below:

Sl.	Document	Genre	Date	Source / Link
1	Pramāṇa-Weighted Intelligence: Reconstructing AI Epistemology and Accountability Through Mīmāṃsā Source-Validity Theory (De and Pal ¹⁹) [A]	Unpublished conference paper, SAAI Conference 2026	2026	The concerned paper has been selected for presentation at the SAAI Conference, 2026, and is under consideration for conference-volume publication, but is not yet published ¹⁹ . Its machine training weight under the PWTA architecture is

				accordingly PWTA(s) = 0.0000, applied by the Āpata lock independently
2	SC Calls for Zero-Tolerance for AI- Generated Fake, Non-Existent or Hallucinated Material During Adjudicatory Process (Parmod Kumar) [B]	Newspaper report, The Statesman (English daily)	02.07.2026	https://www.thestatesman.com/supplements/law/sc-calls-for-zero-tolerance-for-ai-generated-fake-non-existent-or-hallucinated-material-during-adjudicatory-process-1503612441.html
3	The Structural Hole in India's Data Protection Law (Yatharth Chakravarty) [C]	Legal blog article, Live Law	06.06.2026	https://www.livelaw.in/articles/digital-personal-data-protection-act-state-processing-without-consent-proportionality-537036

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

4	Venkatesh Nayak v. Union of India, W.P.(C) No. 177/2026 [C]	Public Interest Litigation petition, Supreme Court of India (unadjudicated)	06.02.2026	https://www.sco-bserver.in/wp-content/uploads/2026/02/Venkatesh-Nayak-v-Union-of-India-Writ-Petition.pdf
5	Zubeen Garg Case: Court Frames Charges Against All Accused, Trial to Begin June 8 [C]	Regional newspaper report on a criminal trial, Pratidin Time (Assam)	26.05.2026	https://www.pratidin-time.com/guwahati-news-breaking-latest/zubeen-garg-case-court-frames-charges-against-all-accused-trial-to-begin-june-8-11875409

Table 3: The five documents subjected to the PWTU-UDE model in the present study. Documents 2 to 5 are cited at notes 24 to 27; the unpublished paper at Sl. 1 is cited at note 19.

All inputs below were derived from direct, first-hand reading and measurement of the complete text of the manuscripts shown above. No inputs were simulated. The raw measurement results are presented in Table 4.

Parameter	SAAI Paper	A: Statesman	B: Live Law	C: PIL	D: Pratidin
-----------	------------	-----------------	-------------	--------	-------------

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com
<https://www.ijalr.in/>

Word count	3,808	1,011	1,476	6,200	242
Paragraphs	92	18	17	98	6
Original tables	4	0	0	1	0
Data numerals	14	0	0	9	0
H(Xs) bits	8.7733	7.537	8.030	9.120	6.820
R_EPS	0.8340	0.8600	0.8550	0.8920	0.7200
D_selection	0.04	0.00	0.15	0.08	0.00
R_CE	1.0000	0.8867	0.9447	1.0000	0.8024
E_s	0.9600	0.7862	0.7586	0.9200	0.6438
C_neo	0.71624	0.00000	0.00000	0.21520	0.00000
S_write	0.9890	0.9643	0.9841	1.000	0.8333
TS_unpub	0.4821	0.0000	0.0000	0.1501	0.0000
PWTA(s)	0.0000	0.0000	0.0000	0.0000	0.0000

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

Table 4: Comparative PWTA-UDE results for the five documents. All values from first-hand measurement and the single-pass formula of Section V. The Āpata lock holds $PWTA(s) = 0.0000$ for every document identically.

All values (in Table 4) were computed using the identical single-pass procedure specified in Section V, using the calibration parameters listed in Table 1. No parameter was altered between the documents. The genre of a document was permitted to inform only D_selection, as shown in Table 5.

Each mathematical parameter of the PWTA-UDE model maps directly to a recognized measurable human trait which can subsequently be compared with a legal analog. Further in Table 5 the concerned analogues, together with remarks explaining the findings obtained on the five documents are presented. The human bias is not missed by the model. It is captured directly and numerically by D_selection, which records the proportion of known contrary positions that a document has chosen to omit. The bias is measured as verifiable omission, without any human opinion about the author or the merits of the argument.

Parameter	Human Trait	Legal Analogue (accepted definition)	Remarks on the Findings
R_EPS	Respect for one's teachers — building upon trusted elders rather than inventing from nothing	“Stare decisis” — to stand by that which has been decided	The PIL is highest at 0.8920, standing on twelve-plus Supreme Court precedents including PUCL (2003), ADR (2002) and both Puttaswamy decisions. Pratidin is lowest at 0.7200, inheriting only from a CID chargesheet

			and a trial-court order. With no apex authority behind it. What a document says is what the document epistemically is.
D_selection	Candour — the willingness to name the argument against oneself	“Suppressio veri” — suppression of a material truth; a breach of uberrima fides (utmost good faith)	This is the human-bias detector of the model. The two reporters (A and D) scored 0.00 indicating that nothing was suppressed, and nothing expected of the genre. The SAAI paper scored 0.04, having cited Searle, Dreyfus and Mohamed against its own thesis voluntarily. The PIL scored 0.08 — adversarial by genre, with the Bench supplying the balance. Live Law scored 0.15: the welfare-delivery

			defense of the DPDP exemption regime is never engaged. Bias is thereby measured as an omission, not as an opinion.
R_CE	Eloquence — the breadth of one's working vocabulary under pressure	“Pari materia” — construed in the same register as the governing law	The SAAI paper and PIL saturate at 1.0000, their entropy exceeding the Tier 1 statutory baseline. Pratidin falls to 0.8024 — a 242-word report of high repetition. It is observed that the blog out-scores the newspaper (0.9447 against 0.8867): analysis ranges wider than reportage, as a structural property of the two genres.
E_s	Integrity — candour and eloquence compounded; neither alone suffices	“Bona fides” — good faith as a composite standard	The squaring of R_CE punishes register drift disproportionately. In the manner in

			which small lapses of good faith taint whole transactions. The Live Law (0.7586) falls below the plainer Statesman (0.7862), while the elegant prose style writing cannot buy back the silence by 0.15.
C_neo	Industriousness — whether the author actually made something, or only commented/ dependent upon what others made	“Best evidence rule” — primary evidence outranks secondary accounts	The great divider of the study. SAAI: 0.71624, from four original tables and fourteen calibration numerals. PIL: 0.21520, from one self-made comparative table of Section 8(1)(j) and nine legally operative figures, including the Schedule penalties of Rs 250, 200, 150 and 50 crore. All three journalism and commentary pieces: 0.00000, as they are

			basically secondary accounts of other persons' primary material.
S_write	Clarity of mind — a sentence whose every word finds its grammatical partner	“Certainty of pleadings” — vague pleadings fail	This is the narrowest spread in the study (0.8333 to 1). All five authors write competently. Style alone ,therefore, separates nothing, which is precisely why the multiplicative AND-rule refuses to permit style to substitute for substance.
M_unobs	Humility before the unverified — trusting less what no institution has yet examined	“Onus probandi” — a burden of proof not yet discharged	A uniform multiplicative discount of 0.85 was applied to all five documents. The noticed absence of an institutional review is itself a valid knowledge (anupalabdhi) as per the mimansha

			philosophy discussed above. The multiplicative form keeps weak documents distinguishable from the weaker ones. The manner in which a verdict of not proven is not the verdict of disproven.
Āpata lock	Prudence — declining to act upon the untested, however promising	“Sub judice” — a matter under judicial consideration is not acted upon	All five documents scored PWTA(s) = 0.0000, as already discussed. No TS_unpub value, however high, purchases a training weight. The audit score informs operators that institutional verification alone feeds machines.
TS_unpub	Earned reputation — the standing accumulated only when every virtue holds at once	“Probative value” — the weight the evidence actually deserves	One zero anywhere zeroes the whole. In the manner in which one fatal defect sinks an otherwise

			flawless case. The final gradient obtained is 0.4821 (SAAI) above 0.1501 (PIL) above 0.0000 (all commentary).
--	--	--	---

Table 5: Correspondence between each PWTA-UDE parameter, its human trait, its accepted legal analogue, and the findings on the five documents. D_selection is the direct, numerical detector of human bias within the model.

6.4 Interpretation of Results

Comparative data revealed a clear structural gradient, as follows:

The unpublished SAAI conference paper is an industrious set with original data, honest disclosure, strong lineage — awaiting only institutional verification that no internal virtue can supply.

The Statesman report is an honest witness that made nothing of its own. It respects the professional commitments of journalism.

The Live Law article appears to be polished and fluent counsel for one side only. Its specific silence on the opposing case costs it more than its fine prose earns.

The PIL petition is a good craftsman among commentators. It can be analyzed that the sub judice principle rightly keeps it out of every training corpus even so, its non-zero score.

However, the Pratidin Time report is a truthful reporter's note — unbiased, unadorned, and epistemically empty.

The resulting scores obtain a sharp gradient, showing 0.4821 for the data-rich academic manuscript, 0.1501 for the professionally drafted petition, and 0.0000 for every form of commentary. The system was designed to draw the discrimination that the extension was designed to make. Also, it was made without any element of human judgment about any author. For every document, the Āpata lock held PWTAs = 0.0000 independently, reaffirming the Stage 3 mathematical formula evolution recorded in Table 2 as discussed above. The consequences of this, gradient for constitutional law, data protection law, and criminal law are examined in Section VII below.

VII. Applications of PWTAs-UDE in the Field of Law

7.1 Constitutional Law and AI-Assisted Judicial Systems

With the deployment of AI systems in Indian judicial research, in relation with the non-arbitrariness principle in *Maneka Gandhi v. Union of India*³ mandates that the AI-assisted judicial processes rest upon principled and traceable epistemic foundations. However, the performance of current AI systems depends on the training it receives which in turn depends on the large data in LLMs, the epistemic positioning of which is mostly unknown. The TS_unpub generated by PWTAs-UDE for a pre-institutional document which is demonstrated in Section VI for the unpublished SAAI Conference paper provides a principled element of exactly such a foundation. A judicial institution or an authority or an AI based software seeking an auditable rationale for why a pre-institutional working or draft document was or was not considered in official reference or in the construction of an AI training corpus can now have a scored, derived, auditable record.

Comparatively, in *Maximillian Schrems v. Data Protection Commissioner*²⁰, the Court of Justice of the EU has established that AI data governance decisions must be subject to independent, verifiable audits. Also, Article 10 of EU AI Act,¹³ mandates that the AI systems deployed in the administration of justice be trained on data of sufficient quality, relevance, and

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

representativeness. But no document till date clearly answers how it can be done. The PWTA-UDE operationalizes these requirements in the Indian context by generating a mathematical approach and an audit-ready score for every pre-institutional source presented to a legal AI training pipeline.

7.2 Data Protection Law and the DPDPA 2023

Section 11 of the DPDPA 2023⁴ confers upon data principals the right to obtain information about how their data have been processed. The karttrva metadata¹ record maintained by the PWTA architecture will log the tier assignment, lineage score, R_EPS, E_s, C_neo, S_write, and TS_unpub for each scored document. This will provide exactly the provenance documentation that enables a data fiduciary/ trustee to respond/ choose / decide to request in respect of an AI training corpus or consideration of the document for official use or for publication as the case may be. In the SAAI paper examined in Section VI: its complete scoring record — PWTA(s) = 0.0000, TS_unpub = 0.4821, with all component values as set out in Table 4 — constitutes precisely this kind of auditable provenance trail. The UK GDPR Articles 13 and 14²¹ and the FTC documentation requirements¹⁵ impose analogous obligations in their respective jurisdictions.

7.3 Ethics of AI and Indian Knowledge Systems

Recommendation of UNESCO, on the Ethics of the Artificial Intelligence, 2021,²² to which India is a signatory, affirms the principles of transparency, fairness, and antibias with valuable human oversight. Also, the IEEE Ethically Aligned Design principles²³ call for traceable and auditable AI decision processes. However, these commitments, without a transparent mathematical computable mechanism for scoring pre-institutional sources, the commitments remain aspirational in practice.

PWTA-UDE suggests an operational commitment at the level of individual document evaluation. The results obtained in Section VI demonstrate this not in principle but in practice. It shows how a real, named, published standard paper has been scored. Through a mathematically supported,

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

first-hand measured, fully traceable procedure, generating a result that a regulator, an editor, or a court can examine, verify, and cite in any proceeding.

7.4 Criminal Law and Administrative Accountability

In the purview of the Bharatiya Nyaya Sanhita, 2023,⁵ professional liability may encompass AI-assisted official decisions. PWTA-UDE's Āpata lock ensures that all unpublished documents, regardless of their TS_unpub score, are categorically excluded from machine training weights. This protection is categorical and not probabilistic. During any challenge to an AI-assisted administrative or criminal law decision, the institution can produce a complete, mathematical, documented account of every document that contributed to its AI system's training. It can demonstrate that pre-institutional material, however high its TS_unpub score, never influenced the model's learned parameters. The results obtained can also be used by officials to determine whether any random document can be used or trusted.

VIII. Conclusion with Suggestions

In the present study, we resolved the unpublished, unreviewed, anonymous document problem in the original PWTA architecture. In this study, the PWTA-UDE extension was developed, specified, and applied to some real manuscripts. The results obtained from the direct measurement of the text yield for SAAI unpublished conference paper, TS_unpub = 0.4821, bounded within [0, 0.85] as required by the architecture's proven boundedness condition. The Āpata lock holds $PWTA(s) = 0.0000$ for all manuscripts this paper independently and categorically, in spite of their TS_unpub score.

The results are the multiplicative-discount form of the formula adopted in Stage 3, after identifying the defects of the earlier additive and subtractive forms produced calibrated, interpretable, granularity-preserving results across the full range of document qualities. The three-stage evolution of the formula, documented in Table 2, suggests a methodological contribution in its own right. It demonstrates concretely why additive combination and flat-constant subtraction are insufficient for epistemic scoring and why the multiplicative-discount

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

form is the only specification consistent with both the mathematical properties required by the architecture. It is well grounded in the philosophical principles of Doṣa-Nirūpaṇa of Mimansha. These findings yield five immediate policy and technical recommendations, which are illustrated below:

First, institutions deploying AI systems in Indian legal, judicial, or administrative contexts ought to adopt PWTa-UDE as the minimum epistemic governance standard for evaluating pre-institutional sources. The framework's computable, fully documented, and easily reproducible design enables compliance with non-arbitrariness. This is also ordained in the Constitution of India under Article 14, the data accuracy obligations under the DPDPA 2023, and the training data quality requirements of the EU AI Act in cross-border AI deployments.

Second, the Government of India, in framing implementing regulations under the DPDPA 2023, ought to prescribe a minimum Corpus Trust Index, as defined in the companion PWTa framework published by De and Pal¹, which can be a qualifying condition. It is intended for use of AI systems, by data fiduciaries/ trustees in contexts that affect individual rights. The mathematical specification of the condition, including the Constraint Optimization procedure for determining the minimum Tier 1 source composition.

Third, the domain-specific calibration procedure for $\tau_{tbl,D}$ and $\tau_{num,D}$, as carried out in the present study for the Indian AI law domain, can be extended to other legal subdomains, such as constitutional law, criminal law, environmental law, and labor law. Thus, the TS_unpub values can be interpreted as domain-specific, percentile-based terms rather than against a single undifferentiated scale.

Fourth, the proposed system can also be used by officials, authorities, academicians, and working law professionals to determine the quality of a manuscript. It is imperative that an automated extraction pipeline, such as a software system that would parse a submitted document, count its paragraph and word totals, identify original un-cited tables, filter qualifying numeric tokens, and compute S_write through the spaCyen_core_web_trf dependency parser, can be constructed and independently validated as a production-grade tool. The concerned pipeline

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

would also permit the formal measurement of the parser's UAS degradation on Indian legal text, which the present study identifies as a limitation and defers to future work. The S_{write} values, are calculated with spaCyen_core_web_trf on all five documents also using the tree-well-formedness criterion. The limitations are discussed in Section IX of this manuscript.

Fifth and finally, a system like Āpata lock principle — the categorical exclusion of unpublished documents from AI machine-training weights, irrespective of their TS_{unpub} score — be codified as a mandatory architectural requirement in any AI governance standard applicable to AI systems in India. Indian law has yet to specify this requirement at the engineering level.

IX. Limitations of the study

The calibration corpus used in this study has 177 documents of mixed domain, the exclusive list of the document is not attached here due to word limit constraints. It was drawn exclusively from English-language sources. The variation of $\tau_{tbl,D}$ and $\tau_{num,D}$ across regional Indian languages and across legal sub-domains other than AI law has not been carried out and is thus identified as future empirical work.

References

1. Moyukh De &Soumavo Pal, 'Pramāṇa-Weighted Intelligence: A Unified Framework For Epistemic Accountability In Neural Architectures Through Mīmāṃsā Source-Validity Theory AndTheAnuvyavasāya Conflict Resolution Protocol' (2026) 3(4) International Journal of Scientific Research and Technology 1214, doi:10.5281/zenodo.19924705.
2. Ashish Vaswani et al., 'Attention Is All You Need' (2017) 31 Advances in Neural Information Processing Systems 5998.
3. Maneka Gandhi v. Union of India, AIR 1978 SC 597.
4. The Digital Personal Data Protection Act, 2023 (Act 22 of 2023), § 11.

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

5. The Bharatiya Nyaya Sanhita, 2023 (Act 45 of 2023), § 15.
6. Ministry of Electronics and Information Technology, Government of India, IndiaAI Mission: Framework Document (2024).
7. NITI Aayog, Responsible AI for All (2021).
8. Timnit Gebru et al., 'Datasheets for Datasets' (2021) 64 Communications of the ACM 86.
9. Margaret Mitchell et al., 'Model Cards for Model Reporting' in Proceedings of the Conference on Fairness, Accountability, and Transparency (2019) 220.
10. National Institute of Standards and Technology, AI Risk Management Framework 1.0 (2023).
11. Parliamentary Standing Committee on Information Technology, Report on Artificial Intelligence: Governance and Accountability (Lok Sabha Secretariat, 2023).
12. Commission Regulation (EU) 2016/679 of 27 April 2016 on the Protection of Natural Persons with Regard to the Processing of Personal Data [2016] OJ L 119/1, art. 5(1)(d) (GDPR).
13. European Parliament and Council of the European Union, Regulation (EU) 2024/1689 on Artificial Intelligence [2024] OJ L 1689, art. 10.
14. Department for Science, Innovation and Technology, United Kingdom, A Pro-Innovation Approach to AI Regulation (White Paper, 2023).
15. Federal Trade Commission (USA), Aiming for Truth, Fairness, and Equity in Your Company's Use of AI (2021).

16. P.M. Sharma, 'A Study of Anupalabdhi as a Pramāṇa in Bhāṭṭa Mīmāṃsā' (1986) 14 Journal of Indian Philosophy 123.
17. Thomas Cover & Joy Thomas, Elements of Information Theory (2d ed., Wiley 2006) 26–29.
18. Matthew Honnibal et al., spaCy: Industrial-Strength Natural Language Processing in Python (2020) <https://spacy.io/models/en#en_core_web_trf> (last visited 13 July 2026), reporting UAS 95.26% for en_core_web_trf on general English benchmarks.
19. Soumavo Pal, 'Pramāṇa-Weighted Intelligence: Reconstructing AI Epistemology and Accountability Through Mīmāṃsā Source-Validity Theory — A Techno-Philosophical Framework for AI Governance from Indian Knowledge Systems' (SAAI Conference Paper, 2026) (unpublished, selected for conference-volume publication; on file with the authors and the SAAI Conference Organising Committee). Cited pursuant to the author's consent; full text available on request.
20. Maximillian Schrems v. Data Protection Commissioner, Case C-362/14, ECLI:EU:C:2015:650 (CJEU, Grand Chamber, 6 October 2015).
21. United Kingdom General Data Protection Regulation 2018 (UK GDPR) (retained EU law), arts. 13–14.
22. UNESCO, Recommendation on the Ethics of Artificial Intelligence, UNESCO Doc. SHS/BIO/PI/2021/1 (2021), adopted by the General Conference at its 41st session.
23. IEEE, Ethically Aligned Design: A Vision for Prioritizing Human Well-being with Autonomous and Intelligent Systems (1st edn, 2019).
24. Parmod Kumar, 'SC Calls for Zero-Tolerance for AI-Generated Fake, Non-Existent or Hallucinated Material During Adjudicatory Process', The Statesman (2 July 2026) <<https://www.thestatesman.com/supplements/law/sc-calls-for-zero-tolerance-for-ai-generated->

For general queries or to submit your research for publication, kindly email us at ijalr.editorial@gmail.com

<https://www.ijalr.in/>

© 2026 International Journal of Advanced Legal Research

fake-non-existent-or-hallucinated-material-during-adjudicatory-process-1503612441.html> (last visited 18 July 2026).

25. Yatharth Chakravarty, 'The Structural Hole in India's Data Protection Law', Live Law (6 June 2026) <<https://www.livelaw.in/articles/digital-personal-data-protection-act-state-processing-without-consent-proportionality-537036>> (last visited 18 July 2026).

26. Venkatesh Nayak v. Union of India, W.P.(C) No. 177/2026 (Supreme Court of India, filed 6 February 2026) (writ petition, unadjudicated) <<https://www.scobserver.in/wp-content/uploads/2026/02/Venkatesh-Nayak-v-Union-of-India-Writ-Petition.pdf>> (last visited 18 July 2026).

27. 'Zubeen Garg Case: Court Frames Charges Against All Accused, Trial to Begin June 8', Pratidin Time (26 May 2026) <<https://www.pratidintime.com/guwahati-news-breaking-latest/zubeen-garg-case-court-frames-charges-against-all-accused-trial-to-begin-june-8-11875409>> (last visited 18 July 2026).