

Analysis of Law School Generative AI Coursework and Exam Policies

Larry Catá Backer (白 轲) (in collaboration with Harvey AI)

W. Richard and Mary Eshelman Faculty Scholar; Professor of Law and International Affairs

Penn State Dickinson Law, a unit of The Pennsylvania State University system | 239 Lewis Katz Building, University Park, PA 16802 | 1.814.863.3640 (direct) || lcb11@psu.edu

Abstract: Structure, Opacity, and Convergence in Law School AI Policies: This report consolidates a multi-stage analysis of generative artificial intelligence (AI) coursework and examination policies across a retrieval-based sample of twelve American law schools and programs. The study demonstrates that law school AI governance cannot be reduced to a single linear spectrum; instead, policies vary independently along three distinct structural axes: default polarity (restrictive versus permissive baselines), drafting style (determinate rules versus interpretive standards), and a two-tier autonomy structure (governing institution-to-instructor and instructor-to-student relationships). Cross-analysis reveals that substantive restrictiveness does not predict structural design, meaning schools with identical baselines often impose vastly different interpretive or administrative burdens on students and faculty. The research identifies a pervasive opacity across the broader legal education sector, noting that a vast majority of ABA-accredited law schools lack retrievable, law-specific public policy texts. This opacity manifests via four distinct patterns: non-existence, disclosed decentralization, active access-gating, and unwritten or oral communication. This lack of public accessibility sits in tension with the fair-notice principles required for academic integrity enforcement. Furthermore, the study finds that law schools rarely author syllabus language independently, relying instead on a small pool of shared template sources. This ecosystem fosters formal convergence on a common taxonomy of policy types while simultaneously permitting wide divergence in substantive local rules. Ultimately, the field develops through a two-tier mechanism: while individual policy documents are structured deductively from a primary principle, the field as a whole evolves inductively and mimetically through horizontal borrowing, imitation, and iterative revisions driven by accumulated institutional experience.

This report consolidates a multi-stage research and analytical effort into a single sequential document. The central finding, developed and refined across successive rounds of research, is that law school policies governing generative AI use in coursework and examinations cannot be reduced to a single restrictive-to-permissive spectrum. Instead, these policies vary along at least three empirically independent structural axes, are documented unevenly across a landscape marked by genuine gaps in public accessibility whose causes are themselves heterogeneous, draw their categorical vocabulary from a small number of widely circulated template sources rather than being independently invented at each institution, and develop through a two-tier process that is internally deductive within any single policy document but inductive and mimetic across the field as a whole. This document is limited to twelve law schools and programs for which retrievable, law-specific coursework or exam policy text was located through iterative web research, against a background of approximately 198 ABA-accredited law schools nationally; it is a non-random, retrieval-based sample and should not be read as a representative survey of American legal education generally. [\[1\]](#) Where evidence is thin, secondary-sourced, or interpretive rather than established, this document says so explicitly rather than smoothing over the distinction.

Scope, Method, and the Twelve-School Dataset

The dataset at the center of this analysis comprises the University of California, Berkeley School of Law; Columbia Law School; the University of Chicago Law School; Stanford Law School, including its Juelsgaard Intellectual Property and Innovation Clinic as a documented sub-institutional case; the University of San Diego School of Law; the Center for Transnational Legal Studies; Fordham University School of Law; Mitchell Hamline School of Law; American University Washington College of Law; the University of Texas at Austin School of Law; Penn State Dickinson Law; and Suffolk University Law School. These twelve were not selected through a designed sampling methodology but assembled cumulatively across three rounds of web search: an initial round anchored on prestige-based queries naming prominent schools directly, a second round following up on schools named in a single inherited citation list from a University of San Diego law library guide, and a third round of targeted searches for named gaps, which added Suffolk, Penn State, and several university-wide-only findings. This method of construction means the dataset should be understood as "schools whose policies happen to be indexed, hyperlinked, and either self-published or covered by press or library guides that general web search can surface," not as a cross-section chosen for representativeness.

Several schools yielded only university-wide, non-law-specific guidance rather than law-school-specific instruments: Northwestern Pritzker School of Law's own library guide page on generative AI was found to be marked, in its own words, "not currently available due to visibility settings"; the University of Michigan, UCLA, and Vanderbilt each produced substantial university-level teaching-center guidance and sample syllabus language without a confirmed law-school-specific counterpart. [\[2\]](#) [\[3\]](#) [\[4\]](#) These are reported separately throughout this analysis rather than folded into the twelve-school primary dataset, because conflating university-wide guidance with a law school's own policy would misrepresent what specifically governs law students. A further set of named schools — Harvard, Yale, NYU, Georgia State, UC Irvine, University of Washington, Washburn, Duke, University of Michigan Law School specifically, UCLA School of Law specifically, and the University of Virginia — were searched but yielded no retrievable law-specific policy text at all, a gap that the analysis in a later section treats as itself a substantive finding rather than a simple absence of data.

Aspects of the methodology and analytical approaches are treated in more detail in Appendix A: sample-selection methodology, definition and treatment of borderline cases, mapping of the observed compliance architectures onto established regulatory-design typologies, incorporation of Penn State's law schools, and a dedicated treatment of faculty autonomy as a cross-cutting structural variable

Three Independent Structural Axes

The core typological contribution of this analysis is a set of three structural axes, derived not from an external normative framework but from open coding of the retrieved primary texts themselves — that is, from categories the schools' own drafters independently used to organize their rules, evidenced by the fact that the same structural divisions (particularly the exam/non-

exam distinction) recur across policies drafted with no apparent coordination among Berkeley, Columbia, Chicago, and Stanford. [\[5\]](#) [\[6\]](#) [\[7\]](#) [\[8\]](#)

Axis One: Default Polarity

Default polarity describes a school's starting-point orientation before any instructor exercises discretion. A **restrictive baseline** treats AI use as prohibited unless an instructor affirmatively loosens the rule; this describes Berkeley, whose policy "forbid[s] the use of AI... to conceptualize, outline, draft, revise, and edit" by default, permitting research only "for the limited purpose of identifying sources", and Columbia, whose "Default Prohibition" bars AI "in (a) any exam or final paper or (b) for aid in drafting any part of work submitted for credit, even if the use is fully documented". [\[9\]](#) [\[10\]](#) The same restrictive-baseline structure appears in Chicago's exam-specific rule, which is "absolutely prohibited... unless the instructor specifies otherwise"; in USD's requirement that "no AGC may be included in written work submitted for credit... unless the instructor explicitly permits such use in writing"; in Fordham's University-wide standard incorporated by its law school, barring use "unless directly authorized by the instructor"; in Mitchell Hamline's plagiarism definition, which treats undocumented AI-generated content as a violation by implication; in Penn State's university-wide exam rule, which states plainly that "students may not use generative AI tools to complete multiple-choice, matching, fill-in the blank, open-ended, or essay exam questions"; and in Suffolk, where a school official is quoted stating that "Suffolk Law has a default policy that prohibits use of generative AI on assignments unless a professor explicitly allows it". [\[5\]](#) [\[11\]](#) [\[12\]](#) [\[13\]](#)

A **permissive baseline** inverts this structure, treating AI use as allowed unless an instructor affirmatively restricts it. The Center for Transnational Legal Studies is the clearest instance: "Unless explicitly prohibited and enforced through supervised assessment conditions, students may reasonably assume AI tool usage is permissible for assignments, provided proper attribution is maintained". Stanford exhibits a genuinely **split polarity**: permissive for idea development and learning support absent a course-specific policy, but restrictive — requiring prior written instructor authorization — for exams and for drafting or revising submitted work. [\[14\]](#) Penn State likewise splits by governance level rather than by use category, combining a restrictive university-wide floor for exams specifically with a complete absence of institutional default for non-exam work, where "the course syllabus and assessment instructions" control entirely. [\[15\]](#) [\[16\]](#) [\[17\]](#) Finally, American University and UT Austin resist confident placement: American University's Academic Integrity Code reportedly contains "no one-size-fits-all" rule at all, while UT Austin's evidence describes a structural exam-format reform — "near-complete elimination of take-home exams" in favor of in-class, software-restricted testing — rather than a permission-based conduct rule comparable to the other schools. [\[18\]](#) [\[19\]](#)

Axis Two: Rule Versus Standard Drafting Style

This axis, drawn from the established rules-versus-standards distinction in regulatory design theory, distinguishes ex ante determinate rules from ex post interpretive standards, independent of substantive restrictiveness. [\[20\]](#) [\[21\]](#) Berkeley is the most **rule-drafted** policy in the dataset, enumerating specific prohibited activities in catalog form — "Asking an AI tool to brainstorm a paper topic or thesis (prohibited conceptualizing)," "Asking an AI tool to propose an

organizational structure for a paper (prohibited outlining)". [22] Columbia and Penn State are similarly rule-drafted, naming covered contexts or exam formats directly rather than relying on an external test. [10] [15] Chicago is the clearest **standard-drafted** policy, asking whether AI use "would constitute academic plagiarism if the generative AI were a human author whose work was used without attribution" — a functional-equivalence test requiring interpretation rather than a determinate catalog. [23] Fordham and Mitchell Hamline are similarly standard-based, folding AI regulation into pre-existing plagiarism doctrine rather than drafting a freestanding enumerated instrument. Stanford, CTLS, and USD occupy a **hybrid** position, combining standard-like general permissions with rule-like bright-line exceptions or documentation duties: Stanford's general guidance is standard-like, but its exam and drafting provisions require a determinate, binary condition — prior written disclosure and authorization — and CTLS's permissive standard is paired with a rule-like mandatory duty to keep "a record of the prompts used to generate content". [14] [24] [25]

Axis Three: The Two-Axis Default/Mandatory Autonomy Structure

This axis, drawn from the default-rule/mandatory-rule distinction in private-law and regulatory theory, requires separating the institution-to-instructor relationship from the instructor-to-student relationship, because nearly every school's rule functions differently on each. [26]

On the **institution-to-instructor** axis, most schools operate a default the instructor may override: Berkeley's instructors may "deviate from the default rule... provided that they do so in writing and with appropriate notice"; Columbia's instructors may move the rule bidirectionally, both loosening and tightening it; Chicago's, Stanford's, USD's, CTLS's, and Fordham's instructors hold comparable, if more often one-directional, override authority. [27] [28] [29] American University and UT Austin show no institutional default to override in the retrieved evidence, while Penn State's university-wide exam rule shows no textual override clause at all, suggesting — tentatively, given only an excerpt was retrieved — a possibly mandatory institutional floor specifically for exams. [15] [30]

On the **instructor-to-student** axis, the evidence shows that whatever rule an instructor sets is, in most schools, mandatory rather than a default the student can access merely through disclosure. Columbia states this most directly: drafting assistance is barred "even if the use is fully documented". [31] Stanford requires not just disclosure but advance authorization: "unless (1) fully disclosed in advance... and (2) explicitly authorized by the instructor in writing prior to the student's use of the tool". [32] USD treats nondisclosure as an aggravating factor rather than treating disclosure as curative, and Chicago's background plagiarism doctrine cannot be waived by attributing work to AI. [33] [34] CTLS is the clear counter-instance: its permissive default is genuinely student-facing, since "students may reasonably assume AI tool usage is permissible... provided proper attribution is maintained", requiring no per-use instructor sign-off.

A bounding floor operates beneath all of this discretion: Stanford states that permitted uses may not "authorize students to contravene standard academic norms concerning plagiarism and accuracy," and that in clinical coursework, "AI use must meet all applicable rules of professional responsibility" — a floor no instructor's permissive policy can waive. [35] [36] A further, previously unremarked governance tier appears at Stanford specifically: its Juelsgaard IP and

Innovation Clinic maintains its own stricter default, distinct from the general Stanford Law policy, stating "as a default, you should not use ChatGPT or any other generative AI tools in completing any of your clinic assignments" — evidence of a program-level tier of discretion between the institution and the individual instructor that the two-tier model, as originally constructed, does not fully capture. [37]

The Central Finding: Axis Independence

Applying these three axes to the twelve-school dataset shows that substantive restrictiveness does not predict drafting style or autonomy structure. Berkeley and Columbia align closely across all three axes — restrictive baseline, rule-drafted, mandatory floor. Chicago matches their restrictive exam polarity but diverges sharply on drafting style, using a standard rather than an enumerated rule. CTLS is the dataset's clearest outlier on baseline polarity and on the instructor-to-student relationship, yet pairs its permissive substance with rule-like procedural documentation demands. Stanford demonstrates that split polarity is a coherent design choice rather than an incomplete one. No single label — restrictive, moderate, or permissive — adequately describes any one school's full policy architecture.

Table 1: Composite Axis Placement Across the Twelve-School Dataset

School	Axis 1: Baseline Polarity	Axis 2: Drafting Style	Axis 3A: Institution→Instructor	Axis 3B: Instructor→Student
UC Berkeley Law	Restrictive [38] [39]	Rule (enumerated catalog) [38]	Default, bidirectional	Mandatory
Columbia Law School	Restrictive [10]	Rule (named contexts) [10]	Default, bidirectional [40]	Mandatory [40]
University of Chicago Law	Restrictive (exam) [5] [41]	Standard (plagiarism-equivalence) [42]	Default [5] [6]	Mandatory [33]
Stanford Law School	Split [27]	Hybrid [29]	Default [43]	Mandatory (prior written authorization) [32]
University of San Diego Law	Restrictive [44]	Hybrid	Default [45]	Mandatory [46]
Center for Transnational Legal Studies	Permissive	Hybrid	Default	Default-like
Fordham Law School	Restrictive [12]	Standard [12] [47]	Default [12]	Not specified

Mitchell Hamline School of Law	Restrictive (implied)	Standard	Not clearly specified	Not specified
American University WCL	Not specified [48]	Standard (per secondary source) [48]	No default identified [48]	Not specified
University of Texas at Austin Law	Not classifiable (structural reform) [49]	Not classifiable	Not specified	Not specified
Penn State Dickinson Law	Restrictive (exam, university-wide); absent (non-exam) [15]	Rule (named exam formats) [15]	Possibly mandatory (exam); absent (non-exam) [15]	Disclosure required where permitted [50]
Suffolk University Law School	Restrictive [51] [52]	Not classifiable	Default [52]	Not specified

Notes: Cells marked "not specified" or "not classifiable" reflect genuine gaps in retrieved evidence rather than inferred values; no cross-school imputation has been performed.

Independence as the Generative Mechanism of Diversity

The diversity observed across the twelve-school dataset is not simply the additive result of three separate disagreements — schools disagreeing about polarity, disagreeing about drafting style, and disagreeing about autonomy structure, as though these were three unrelated votes being tallied. Rather, the diversity is combinatorial: because the three axes are empirically independent of one another, a school's position on one axis constrains almost nothing about where it will land on the other two, and the effective policy space is therefore far larger than the small number of category values on any single axis would suggest. With even a modest three-to-four value range on each axis, the theoretical combinatorial space runs into the dozens of distinct architectures, and the twelve schools in this dataset already populate a meaningful fraction of that space with genuinely distinct profiles rather than clustering into one or two dominant types. This is the central reason a single restrictive-to-permissive label cannot describe any school adequately: restrictiveness is a property of only one axis, and the other two axes determine how that restrictiveness is expressed, enforced, and experienced by instructors and students.

Polarity and Drafting Style: Restrictiveness Delivered Through Different Interpretive Mechanisms

The interaction between default polarity and drafting style determines not *how much* AI use is permitted but *how the boundary of permission is discovered* in a given case. Berkeley and Chicago both adopt a restrictive baseline for exams — AI use is prohibited by default in both — yet they arrive at that restriction through opposite drafting mechanisms. [23] [83] Berkeley's restriction is rule-drafted: it enumerates specific prohibited activities, such as "asking an AI tool to compose a paragraph summarizing a legal rule for use in a paper", meaning that determining whether a given student action violates the policy is, in principle, a matter of matching conduct against an explicit list. [3] [4] Chicago's restriction is standard-drafted: the operative test is whether the conduct "would constitute academic plagiarism if the generative AI were a human author whose work was used without attribution", meaning that determining a violation requires an evaluator to reason by analogy to a separate, pre-existing body of doctrine. [5]

The practical consequence is that two schools with identically restrictive substantive baselines impose very different burdens on the parties who must apply the rule. Berkeley's catalog can, in an easy case, be applied by a student or instructor without any need to consult external doctrine; Chicago's standard requires active interpretive judgment in every case, shifting adjudicative cost from the drafting stage (where Berkeley invested effort enumerating instances) to the application stage (where Chicago's evaluators must reason case by case). This means "restrictive" schools are not equally predictable or equally administrable, and the drafting-style axis is what determines which of these two very different experiences a restrictive rule produces.

The same interaction runs in the other direction for permissive baselines. CTLS's permissive default — "students may reasonably assume AI tool usage is permissible for assignments, provided proper attribution is maintained" — is standard-like in its general orientation, granting broad latitude subject to a loosely specified attribution norm. But CTLS pairs that permissive standard with a rule-like, determinate documentation requirement: a mandatory "record of the prompts used to generate content". [6] The interaction here produces a hybrid experience distinct from either a purely permissive-and-standard-based policy or a purely permissive-and-rule-based one: substantive freedom is wide, but the administrative burden of demonstrating compliance is specific and exacting. A school that combined permissive polarity with a purely standard-based compliance mechanism (say, only "students should be transparent about AI use" with no further specification) would produce meaningfully more discretion, and less predictability in enforcement, than CTLS's actual combination.

Polarity and Autonomy Structure: Where the Substantive Default Actually Operates

The interaction between default polarity and the institution-to-instructor autonomy relationship determines where in the institutional hierarchy substantive disagreement is actually permitted to occur, which in turn shapes how much real-world variation exists beneath a school's nominal single policy. Because Berkeley, Columbia, Chicago, Stanford, USD, and Fordham all pair a

substantive default (restrictive, in most of these cases, or split in Stanford's) with instructor override authority, the school-level policy label describes only the *fallback* condition, not the actual lived experience in any classroom where an instructor has exercised override authority. [7] [8] This suggests that within a single school with a nominally restrictive default, the true range of experienced policy can span from fully restrictive (where an instructor never states an alternative) to fully permissive (where an instructor has affirmatively adopted a more permissive syllabus statement), and the twelve-school comparison in the earlier analysis, which necessarily reports only the institutional default, systematically understates the true diversity present *within* each school's course offerings. CTLS's permissive default with instructor-restriction authority produces the mirror image: the school-level label is permissive, but any individual course could be more restrictive if its instructor has exercised the override.

This interaction becomes especially consequential where the autonomy structure shows no institutional default at all, as at American University and UT Austin. [9] [10] In these cases, the polarity axis is not merely unknown at the school level; it is, in a meaningful sense, not a school-level property at all, since there is no institutional fallback for the polarity axis to describe. Diversity at these schools is therefore not bounded by an institutional default the way it is at Berkeley or CTLS; it is bounded only by whatever variation exists among individual instructors' independently originated policies, which — absent any institutional floor — could in principle range across the entire spectrum without any single school-level statement constraining that range. The interaction of "no default" on the autonomy axis effectively removes the polarity axis's descriptive power entirely for these schools, which is why they were marked "not specified" rather than assigned a value in the earlier table: the axes are not just independent in general, but can become mutually inapplicable to one another under certain combinations.

Drafting Style and Autonomy Structure: Predictability of Override and the Cost of Deviation

The interaction between drafting style and the autonomy structure governs how easy or difficult it is for an instructor to actually exercise the override authority the autonomy axis grants. Where a school's default is rule-drafted, as at Berkeley, an instructor wishing to deviate has a clear, enumerated baseline to deviate *from* — the instructor can specify precisely which items on Berkeley's catalog of prohibited activities are being relaxed for their course, and the deviation itself can be drafted with the same rule-like precision, since the underlying default was already rule-like. [11] Where a school's default is standard-drafted, as at Chicago, an instructor's deviation is harder to specify with equivalent precision, because the baseline itself is a general test rather than an enumerated list; an instructor who wants to permit some AI use beyond Chicago's plagiarism-equivalence default must draft language that itself either introduces a new standard or awkwardly grafts specific rule-like exceptions onto a standard-based baseline. This suggests that the drafting-style axis has a downstream effect on the *quality and clarity* of instructor-level deviations even though the autonomy-structure axis is what formally authorizes those deviations in the first place — a rule-drafted institutional default makes instructor override easier to draft cleanly, while a standard-drafted default makes instructor override more likely to introduce its own interpretive ambiguity.

The interaction also affects the instructor-to-student sub-axis in a related way. Stanford's combination of a hybrid drafting style with a mandatory instructor-to-student floor produces an unusually demanding compliance path: because the exam and drafting exceptions are rule-like (a determinate, binary "prior written authorization" requirement) layered on top of a standard-like general permission, a student cannot satisfy the mandatory floor through good-faith disclosure alone, as could happen under a purely standard-based mandatory floor like Chicago's plagiarism-equivalence test, where at least in principle a student's honest, well-attributed use might satisfy the underlying norm even without a separate authorization step. [12] Stanford's rule-like procedural exception makes the mandatory floor stricter in practice than its permissive general orientation would suggest, precisely because of how the drafting-style axis interacts with the instructor-to-student autonomy axis at the specific junctures (exams, drafting) where Stanford chose rule-like rather than standard-like specification.

Table: Illustrative Combinatorial Profiles Among the Twelve Schools

School	Polarity × Drafting Interaction	Polarity × Autonomy Interaction	Drafting × Autonomy Interaction
UC Berkeley Law	Restrictive baseline delivered via enumerated catalog: easy-case predictability [13]	Restrictive default with bidirectional override: true classroom range spans the full spectrum	Rule-like default makes instructor deviations easy to draft with matching precision [11]
University of Chicago Law	Restrictive baseline delivered via plagiarism-equivalence standard: case-by-case interpretive burden [14]	Restrictive default, instructor may specify otherwise: classroom range varies, but baseline interpretation itself is unstable [1] [15] [16]	Standard-like default makes instructor deviations harder to draft with equivalent precision
Center for Transnational Legal Studies	Permissive baseline delivered via standard-like general permission plus rule-like documentation duty: wide substantive freedom, exacting compliance burden [17]	Permissive default, instructor may restrict: classroom range runs from fully permissive to instructor-restricted	Hybrid default's rule-like documentation duty persists regardless of instructor's substantive choice [18] [19]
Stanford Law School	Split polarity delivered via hybrid drafting: permissive zones are standard-like, restrictive zones	Default with instructor discretion, but exam/drafting exceptions carry a mandatory, non-	Rule-like exceptions make the mandatory floor stricter in practice than the

	(exam/drafting) are rule-like [12]	waivable authorization step [12]	general permissive orientation suggests
American University WCL	No institutional polarity to combine with drafting style; case-by-case per secondary source [10] [20]	No default exists to override; instructor originates rule directly	Not classifiable given absence of institutional baseline [20]

Notes: This table illustrates interaction effects for a representative subset of the twelve schools rather than exhaustively cataloguing all combinations; cells reflect the interpretive synthesis of previously cited primary sources rather than new evidentiary claims.

Why the Interaction, Rather Than Any Single Axis, Explains the Dataset's

Diversity

The cumulative effect of these pairwise interactions is that the twelve-school dataset exhibits more genuine diversity than a simple count of category values on any one axis would predict, because each school's full policy identity is a *product* of its three axis positions rather than a sum. Two schools sharing a restrictive polarity (Berkeley and Chicago) can still produce meaningfully different compliance experiences once drafting style is factored in; two schools sharing a permissive polarity and a similar autonomy structure (CTLs, and Stanford's non-exam default) can still diverge once the specific procedural demands attached to their exceptions are considered. This combinatorial structure is also what makes the earlier finding of axis independence more than a technical curiosity: it means that any attempt to reform or harmonize law school AI policy by focusing on a single axis — for instance, a national push toward uniform substantive restrictiveness — would leave the other two axes, and therefore a substantial share of the practical diversity in student and instructor experience, entirely untouched. The diversity observed across this dataset is, in this sense, structurally overdetermined: it would persist even if every school in the sample converged on identical substantive polarity, because drafting style and autonomy structure remain free to vary independently and would continue generating distinct policy architectures on their own.

The Transparency and Opacity Question

A distinct line of inquiry examined why so many named law schools yielded no retrievable policy text, and whether that absence reflects a single phenomenon or several. Four analytically distinct patterns emerged, each requiring different evidence to confirm, and conflating them would overstate or understate the underlying concern depending on direction.

Non-existence means no rule exists at any organizational level, written or oral. This cannot be confirmed for any specific named law school in this dataset, but is supported at the sector-wide level: UNESCO's May 2023 survey of over 450 institutions found that fewer than 10% had any

formal policy or guidance on generative AI use, with universities somewhat more likely than schools to have guidance (roughly 13% versus 7%). [\[53\]](#) [\[54\]](#) [\[55\]](#)

Disclosed decentralization means an institution openly states that authority is delegated, with no further access barrier once that statement is found. American University's Academic Integrity Code fits this pattern, since the retrieved description states plainly there is "no one-size-fits-all" rule. [\[30\]](#) Vanderbilt's university-wide guidance is an even cleaner instance, specifying precisely what governs when an instructor is silent: "the university permits students to use generative AI tools, but they must disclose all generative AI usage". [\[56\]](#)

Access-gating means content was authored and then affirmatively restricted. The strongest confirmed instance in this entire research effort is Northwestern Pritzker School of Law's own library guide page, which displays the message "This page is not currently available due to visibility settings" — direct, first-party evidence of authored content subsequently made inaccessible. [\[4\]](#) Fordham's hosting of its academic regulations via a Google Drive PDF link and Suffolk's reference to an unretrieved "Academic Rules and Regulations" document are consistent with, but do not confirm, this same pattern; they remain properly classified as ambiguous rather than proven instances of gating. [\[57\]](#) [\[58\]](#)

A fourth, residual pattern — **unwritten or oral policy** — is distinct from all three: the rule is operative and known within the institution, but no document exists for anyone, inside or outside the institution, to find. The UNESCO survey found that among institutions reporting any guidance at all, roughly 40% said that guidance had only ever been communicated orally, and close to 20% of respondents did not know whether their own institution had a policy at all — evidence that this pattern is prevalent and that it blurs the line between non-existence and inaccessibility, since a policy unknown even to affiliated respondents is functionally opaque regardless of whether a document technically exists. [\[59\]](#) [\[60\]](#)

Table 2: The Four Patterns, Their Tests, and Confirmed Instances

Category	Defining Test	Confirmed Instance	Confidence
Non-existence	No rule exists at any level, written or oral	Not confirmed for any named law school; supported at sector-wide survey level [59] [60] [61]	Sector-wide: moderate-high; school-specific: unconfirmed
Disclosed decentralization	Publicly findable statement affirmatively delegates rule-making	American University WCL; Vanderbilt [2] [30]	High
Access-gating	Content authored, then placed behind a restriction	Northwestern Pritzker School of Law [62]	High (this one instance)

Unwritten/oral policy	Operative but undocumented	Not confirmed for any named law school; established as prevalent sector-wide (~40%) [59]	Sector-wide: moderate-high; school-specific: unconfirmed
-----------------------	----------------------------	--	--

Credibility Implications and the Hypothesized Meta-Policy

This opacity sits in tension with universities' and law schools' own frequent public self-presentation as champions of open inquiry and reasoned, published rules — a tension sharpened by the doctrinal fact that academic integrity enforcement rests on a fair-notice principle explicit in these same schools' own texts, such as Berkeley's requirement that instructor deviations be made "in writing and with appropriate notice". A 2025 Higher Education Quarterly study of 124 undergraduates and seven faculty members at two U.S. R1 universities found that "students reported that while classroom-level policies are recognised, institutional guidelines remain unclear" — evidence that the fair-notice precondition these institutions rely on to justify sanctioning students may not reliably hold even for the students the rule is meant to bind. [\[63\]](#) The same study found that "students aware of AI policies are less likely to use AI for writing and research", establishing that policy awareness has a measurable behavioral consequence, which sharpens the fairness stakes of any awareness gap. [\[64\]](#)

One commentary — explicitly interpretive rather than empirical, and clearly so labeled here — argues the slow, uneven institutional response "reflected a deliberate, if unspoken, institutional choice," because addressing the problem properly "would require spending money, setting clear grading standards and risking conflict with students". [\[65\]](#) Testing this argument's structure against the assembled evidence, three recurring features are *consistent with* (though not proof of) a cost-avoidance and discretion-preserving logic: the near-universal default-with-instructor-override architecture, which disperses the operative rule content away from any single centrally catalogued document; the roughly 40% rate of oral-only policy communication, which avoids creating a citable document that could later be invoked against the institution; and the explicit caution, found at CTLS and Vanderbilt, against relying on AI-detection scores as determinative evidence, indicating institutions are already managing evidentiary exposure carefully in adjacent respects. [\[66\]](#) [\[67\]](#)

If such a latent meta-policy exists, its hypothesized logic — labeled here explicitly as **unproven hypothesis, not established fact** — would prefer defaults over mandates, dispersed syllabus-level disclosure over centralized publication, and open-ended, revisable language (Chicago's own policy states it "may change as circumstances warrant") over settled commitments. [\[68\]](#) This entire pattern remains equally explicable by mundane, non-strategic causes: genuine technological uncertainty, resource constraints, and ordinary institutional lag between adopting and publishing a norm. This research cannot distinguish between these explanations, because only the external pattern, not internal institutional deliberation, was accessible to it.

Model and Template Syllabus Language, and Discursive Intersubjectivity

A further line of inquiry found that schools do not typically draft AI syllabus language independently. USD's law library curates five named, selectable archetypes — "Prohibited Use," "Mandatory Disclosure," "Permissible vs. Non-permissible uses," "Encouraging Use," and "Prior Consultation". [69] [70] Stanford's Robert Crown Law Library maintains a parallel "Syllabus Statements for Generative AI Usage" resource. [71] Penn State's library guide links directly to Lance Eaton's widely circulated academic spreadsheet and to a Harvard Bok Center illustrated rubric. [72] Michigan's central AI resource site reproduces categorized examples attributed to external sources including Temple University and the multi-institutional Sentient Syllabus Project. [73] [74] UCLA states its own sample language is "adapted from UIC's AI Writing Tools guide", naming the donor institution directly.

Three patterns emerge from this template ecosystem.

Hub concentration: a small number of source documents — Eaton's spreadsheet, the Sentient Syllabus Project, UIC's guide, the AALS 2023 interview on Berkeley's policy — recur as reference points across otherwise unconnected schools.

Convergent categorization despite non-identical wording: USD's five categories, Michigan's three, and UCLA's three are not verbatim copies of each other, yet each carves the same underlying decision space into a similarly small number of discrete types, indicating a shared conceptual ontology has stabilized across the field independent of exact phrasing.

Explicit cross-institutional citation as a legitimating move: UCLA's direct attribution to UIC, and Stanford's clinic citing the AALS interview specifically discussing Berkeley's policy, show institutional authority being transmitted horizontally between peer schools rather than independently justified from first principles at each site. [75]

This patterning suggests some support for characterizing the field as discursively intersubjective: a shared vocabulary lets one school's policy remain legible to another's faculty and students without each institution reconstructing the underlying logic from scratch, even as substantive outcomes — restrictive versus permissive — remain genuinely unsettled and divergent.

The Mechanism: A Shared Grammar, Not a Shared Text

The apparent paradox — that law schools draw from the same small pool of template sources yet arrive at substantively different, often conflicting rules — dissolves once the object being borrowed is properly identified. What circulates across institutions is not policy content in the sense of a specific permission or prohibition, but a **categorical scaffolding**: a menu of archetypal policy *types* that a drafter can select among and then fill with locally determined substance. This distinction between borrowing a form and borrowing a conclusion is the key to understanding why the same source material generates convergence on structure and divergence on outcome simultaneously, rather than one or the other.

How Shared Sources Produce Convergence

The convergence effect operates through what might be called an ontology-stabilizing function. When USD's law library presents five discrete archetypes — "Prohibited Use," "Mandatory Disclosure," "Permissible vs. Non-permissible uses," "Encouraging Use," and "Prior Consultation" — it is not merely offering sample text; it is proposing a taxonomy of the *kinds* of positions a policy can occupy. [1] [2] Once other institutions encounter this same taxonomy, whether directly or through independently arriving at a similar decomposition, a common vocabulary becomes available for describing what any given school's policy *is*, even before anyone has decided what that policy should *say*.

This may explain why Michigan's three-category scheme (encouraged, allowed-with-distinction, misconduct) and UCLA's three-category scheme (not permitted, limited use with citation, permitted with no restrictions) are not verbatim copies of USD's five categories, yet all three partition essentially the same decision space using a comparably small number of discrete bins. [3] Convergence, in other words, occurs at the level of the *classification system* itself — the shared sense of what dimensions matter (is drafting covered? is disclosure required? is the exam context distinct from the paper context?) — rather than at the level of the specific rule chosen within that system.

This convergence is reinforced by explicit citation practices that function as legitimating moves. UCLA states its sample language is "adapted from UIC's AI Writing Tools guide", and Stanford's Juelsgaard Clinic resource page cites the AALS interview specifically discussing the creation of Berkeley Law's policy. [4] [5] Each act of citation transmits not just wording but the *authority* of the borrowed structure — invoking a peer institution's credibility to validate the categorical choices being made locally. Over enough iterations of this kind of horizontal citation, the field converges on a stable, mutually intelligible way of talking about AI policy even across schools that have never directly coordinated, producing the discursive intersubjectivity discussed in the prior analysis: any law school's policy becomes legible to another school's faculty or students because both are drawing on the same underlying grammar of policy types.

How the Same Sources Simultaneously Produce Divergence

Divergence arises precisely because these template sources are explicitly designed and presented as **selectable menus rather than mandates**. USD's library guide does not tell an instructor which of its five archetypes to adopt; it curates the range of options and leaves the substantive choice — restrictive, permissive, disclosure-based — entirely to the local drafter. [6] The same is true of NYU Steinhardt's automated syllabus-statement generator, which produces sample language "based on your unique course information and preferences" — the tool standardizes the categorical shape of the output while leaving its substantive content to vary by user input. [7]

This may suggest that the very feature that produces formal convergence (a shared, limited set of archetypes) is structurally what enables substantive divergence: because the archetype is presented as a template to be filled rather than a rule to be adopted wholesale, each institution's local values, risk tolerance, disciplinary culture, and institutional history determine which slot in the shared taxonomy it occupies, and different institutions predictably choose different slots.

The twelve-school dataset demonstrates this divergence concretely even among schools using structurally similar drafting approaches. Berkeley and Columbia both adopt a rule-drafted, restrictive-baseline structure, yet Berkeley's substantive content bars conceptualizing, outlining, and drafting outright by default, while Columbia's substantive content, though similarly restrictive, is phrased around named contexts (exams, final papers) rather than named cognitive activities — a variation in what counts as covered conduct even within the same formal category. [8] [9] [10] More strikingly, CTLS occupies the opposite pole of the same shared taxonomy: its permissive-baseline structure is drawn from the identical menu of possible policy types (default-plus-override, disclosure requirement, plagiarism-adjacent enforcement) that produces Berkeley's restrictive outcome, but CTLS fills the permissive slot rather than the restrictive one. The shared grammar did not push these schools toward the same substantive conclusion; it simply gave them a common language in which to express conclusions that remain genuinely opposed.

Why This Distinction May Matter Analytically

This convergence-divergence pattern is itself diagnostic of an inductive, mimetic field rather than a deductive, axiomatic one, consistent with the developmental-logic finding discussed previously. If the field were deductive — if a single governing principle logically generated each school's specific rule — one would expect convergence on *substance*, since a shared premise applied validly should yield a shared conclusion. What is observed instead is convergence on *form* paired with persistent divergence on *substance*, which is the signature of institutions observing and borrowing each other's classificatory tools while independently exercising local judgment about where within those tools to land. The template sources function less like a constitution being applied and more like a shared drafting vocabulary being adapted — comparable to how common-law jurisdictions might converge on the same doctrinal categories (duty, breach, causation, damages) while reaching opposite outcomes on any given set of facts. The templates lower the cost of drafting and make disparate policies comparable to one another, but they do not, and were never designed to, resolve the underlying substantive disagreement about how restrictive or permissive AI use in legal education ought to be. That disagreement remains exactly where it was before the templates existed; only the vocabulary for expressing and comparing it has been standardized.

Inductive and Deductive Structure: A Two-Tier Finding

The final analytical question addressed whether the field's overall development is inductive, grounded in mimetic iteration, or deductive, with a governing principle preceding and generating specific applications. The evidence supports neither label applied uniformly, but rather a two-tier structure that separates the internal logic of individual documents from the logic of the field's development over time.

Within a single school's policy document, the structure is frequently deductive: a governing principle is stated first, and specific rules are derived as applications of it. Berkeley states its purpose — equipping students "to perform activities constitutive of excellent lawyering" — before enumerating specific prohibited activities that follow from it. [76] Chicago states a general functional-equivalence test before deriving specific permitted and prohibited outcomes

from it. [42] [77] [78] [79] In both cases, a new, unanticipated use of AI is, in principle, resolvable by applying the stated standard, without requiring a new enumerated example — the defining feature of a deductive system.

Across the field as a whole, however, the evidence points to an inductive, mimetic-iterative process. No single, field-wide governing principle precedes or generates the individual schools' policies. The timeline — an initial adoption wave clustered in Autumn 2023, followed by a substantial revision wave in 2025–2026 across Columbia, Chicago, Berkeley, CTLS, and UT Austin — may be usefully explained as iterative adjustment to observed experience than as the unfolding of an antecedent axiom. Berkeley's own policy moved from comparatively permissive in 2023, when AI could be used "to perform research in ways similar to search engines... and for other functions attendant to completing an assignment", to substantially more restrictive by 2026, barring conceptualizing, outlining, and drafting outright — a trajectory driven by accumulated institutional experience rather than one stable underlying principle. [38] [80] UT Austin's 2026 reform can be understood as explicitly framed as a response to observed practice, warning that classroom time may be "the sole context in which a professor can be certain" a student is learning unaided. [81] [82] The template-circulation evidence reinforces this: the field shows convergence on categorical *form* (a shared menu of three-to-five policy archetypes) alongside persistent divergence on substantive *outcome* (restrictive versus permissive baselines) — precisely the signature of a precedent-and-imitation process rather than an axiomatic one.

What may be a significantly defensible overall characterization is therefore a two-tier structure, analogous to common-law doctrinal reasoning: any single judicial opinion may present a deductive syllogism of rule, application, and holding, while the body of doctrine as a whole develops through accretion, borrowing, and revision in response to observed consequences rather than derivation from one unifying antecedent principle. This is offered strictly as a structural description, not as a normative judgment about whether such a development pattern is desirable or well-suited to governing a fast-moving technology.

Conclusion

Taken together, the layers of this analysis converge on several bounded conclusions, each qualified by the limits of the evidence available. Substantively, the twelve schools in this dataset cannot be sorted along a single restrictive-to-permissive spectrum; three independent structural axes — default polarity, drafting style, and a two-part autonomy structure — are needed to describe any one school's architecture, and schools that align on one axis frequently diverge on another. Procedurally, the broader landscape in which these twelve schools sit is marked by a documented and multi-causal opacity problem: some absence of retrievable policy reflects genuine non-existence, some reflects openly disclosed decentralization that is not itself concealed, some reflects at least one confirmed instance of a school affirmatively restricting access to previously authored guidance, and a substantial share reflects a sector-wide pattern of unwritten, orally transmitted rules that this research's methods cannot fully penetrate. This opacity sits in real tension with the transparency values these institutions publicly espouse, and while one interpretive account frames that tension as a deliberate institutional strategy, the

available evidence supports only the existence and behavioral consequences of the opacity itself, not a confirmed account of institutional motive.

Discursively, the field's policy language is not independently invented at each school but drawn from a small set of recurring template sources, producing convergence on categorical form even amid continued substantive disagreement. And developmentally, the field exhibits a two-tier structure in which individual policies are often internally deductive while the field as a whole evolves inductively, through observation, borrowing, and iterative revision rather than derivation from a stable, antecedent principle. Any reader seeking to extend this analysis to the roughly 186 ABA-accredited law schools not examined here, or to establish institutional motive behind the documented opacity with greater confidence, would need direct institutional outreach or internal documentation of a kind not accessible through the web-search methods used throughout this research.

References

1. [ABA-Approved Law Schools](#)
2. [Academic Integrity and Generative AI | Generative AI | Vanderbilt University](#)
3. [Academic Integrity and Generative AI | Generative AI | Vanderbilt University](#)
4. [Generative AI in Law School and Higher Education - Generative AI Resources - Pritzker Legal Research Center at Northwestern Pritzker School of Law](#)
5. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
6. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
7. [Artificial Intelligence Policy - UC Berkeley Law](#)
8. [Artificial Intelligence Policy - UC Berkeley Law](#)
9. [Artificial Intelligence Policy - UC Berkeley Law](#)
10. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
11. [Standards of Academic Integrity - Fordham University](#)
12. [Standards of Academic Integrity - Fordham University](#)
13. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
14. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
15. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
16. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
17. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
18. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
19. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
20. [\[PDF\] Rules vs. Standards in Private Ordering](#)
21. [\[PDF\] Rules vs. Standards in Private Ordering](#)
22. [Artificial Intelligence Policy - UC Berkeley Law](#)
23. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
24. [Policy on the Use of Artificial Intelligence in Academic Work](#)
25. [Policy on the Use of Artificial Intelligence in Academic Work](#)
26. [A Theory of Mandatory Rules: Typology, Policy, and Design | Texas Law Review](#)
27. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
28. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
29. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
30. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
31. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
32. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
33. [4.9 Law School Policy on Generative AI](#)
34. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
35. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
36. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
37. [Clinic, Law School, and University AI Policies and Syllabus Language - Juelsgaard Intellectual Property and Innovation Clinic - Stanford Law School](#)

38. [Artificial Intelligence Policy - UC Berkeley Law](#)
39. [Artificial Intelligence Policy - UC Berkeley Law](#)
40. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
41. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
42. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
43. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
44. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
45. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
46. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
47. [Standards of Academic Integrity - Fordham University](#)
48. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
49. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
50. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
51. [AI enters the classroom as law schools prep students for a tech-driven practice](#)
52. [AI enters the classroom as law schools prep students for a tech-driven practice](#)
53. [UNESCO survey: Less than 10% of schools and universities have formal](#)
54. [UNESCO survey: Less than 10% of schools and universities have formal](#)
55. [UNESCO survey: Less than 10% of schools and universities have formal](#)
56. [Academic Integrity and Generative AI | Generative AI | Vanderbilt University](#)
57. [Artificial Intelligence \(AI\) - Law Faculty Guide to Artificial Intelligence - LibGuides at Moakley Law Library, Suffolk University](#)
58. [Academic Regulations | Fordham School of Law](#)
59. [UNESCO survey: Less than 10% of schools and universities have formal](#)
60. [UNESCO survey: Less than 10% of schools and universities have formal](#)
61. [UNESCO survey: Less than 10% of schools and universities have formal](#)
62. [Generative AI in Law School and Higher Education - Generative AI Resources - Pritzker Legal Research Center at Northwestern Pritzker School of Law](#)
63. [Exploring the Effectiveness of Institutional Policies and Regulations for Generative AI Usage in Higher Education - Jiang - 2025 - Higher Education Quarterly - Wiley Online Library](#)
64. [Exploring the Effectiveness of Institutional Policies and Regulations for Generative AI Usage in Higher Education - Jiang - 2025 - Higher Education Quarterly - Wiley Online Library](#)
65. [AI is degrading the value of a university degree - LSE Impact](#)
66. [Policy on the Use of Artificial Intelligence in Academic Work](#)
67. [Academic Integrity and Generative AI | Generative AI | Vanderbilt University](#)
68. [4.9 Law School Policy on Generative AI](#)
69. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
70. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)

71. [Clinic, Law School, and University AI Policies and Syllabus Language - Juelsgaard Intellectual Property and Innovation Clinic - Stanford Law School](#)
72. [Getting Started - Generative AI for Legal Educators - Dickinson Law Library at Penn State Dickinson Law](#)
73. [Course Policies & Syllabi Statements - U-M Generative AI](#)
74. [Course Policies & Syllabi Statements - U-M Generative AI](#)
75. [Clinic, Law School, and University AI Policies and Syllabus Language - Juelsgaard Intellectual Property and Innovation Clinic - Stanford Law School](#)
76. [Artificial Intelligence Policy - UC Berkeley Law](#)
77. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
78. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
79. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
80. [Faculty Perspectives: Does Your Law School Need a Policy on Generative Artificial Intelligence? • AALS](#)
81. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
82. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
83. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
84. [Artificial Intelligence Policy - UC Berkeley Law](#)
85. [Artificial Intelligence Policy - UC Berkeley Law](#)
86. [Artificial Intelligence Policy - UC Berkeley Law](#)
87. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
88. [Policy on the Use of Artificial Intelligence in Academic Work](#)
89. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
90. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
91. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
92. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
93. [Artificial Intelligence Policy - UC Berkeley Law](#)
94. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
95. [Artificial Intelligence Policy - UC Berkeley Law](#)
96. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
97. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
98. [4.9 Law School Policy on Generative AI | University of Chicago Law School](#)
99. [Policy on the Use of Artificial Intelligence in Academic Work](#)
100. [Policy on the Use of Artificial Intelligence in Academic Work](#)
101. [Policy on the Use of Artificial Intelligence in Academic Work](#)
102. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
103. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
104. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
105. [Course Policies & Syllabi Statements - U-M Generative AI](#)

106. [Creating a Generative AI Policy for Your Course - UCLA Teaching & Learning Center](#)
107. [Clinic, Law School, and University AI Policies and Syllabus Language - Juelsgaard Intellectual Property and Innovation Clinic - Stanford Law School](#)
108. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
109. [Academic Integrity and Syllabus Support - NYU Steinhardt](#)
110. [Artificial Intelligence Policy - UC Berkeley Law](#)
111. [Artificial Intelligence Policy - UC Berkeley Law](#)
112. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
- 113.

APPENDIX A Methodology and Analytical Issues

This Appendix addresses in more detail five aspects of the study: sample-selection methodology, definition and treatment of borderline cases, mapping of the observed compliance architectures onto established regulatory-design typologies, incorporation of Penn State's law schools, and a dedicated treatment of faculty autonomy as a cross-cutting structural variable.

(1) How the Ten-School Sample Was Selected

The sample was not constructed through a designed sampling methodology (random, stratified, or purposive in the formal social-science sense). It is more accurately described as a **cumulative snowball/convenience sample** built across three successive rounds of web search, where each round's query terms determined what could be found, and the ten schools that ultimately yielded retrievable law-school-specific text are simply the schools for which such text surfaced in the search index and was successfully extracted. Four distinct mechanisms fed the sample:

1. **Prestige-anchored initial queries.** The first search round used queries naming prominent schools directly (e.g., "top law schools Harvard Yale Stanford Columbia NYU"), which by construction privileges institutions with high search-engine visibility and media coverage [prior turn methodology].
2. **Inherited citation lists from a single secondary source.** Several schools in the sample (Fordham, Mitchell Hamline, Georgia State, UC Irvine, University of Washington, Washburn) entered the corpus not through independent discovery but because a University of San Diego law library guide's "Other University Policies" section named them. [1] [2] Only Fordham and Mitchell Hamline among that group actually yielded retrievable substantive text on follow-up search; the rest remained citation-only.
3. **Follow-up targeted queries for named gaps.** A second round specifically searched for schools identified but not yet retrieved (Columbia, NYU, Stanford, Duke, Michigan), which added Columbia and Stanford to the retrievable set.
4. **News-driven discovery.** UT Austin entered the sample not through a policy-document search at all, but through a Bloomberg Law news article about a dean's memo — meaning its inclusion reflects that a journalist covered it, not that its policy documentation is more complete or more representative than that of unreported schools.

[3]

The consequence is that the ten-school sample should be understood as **"schools whose AI policies happen to be indexed, hyperlinked, and either self-published or covered by press or library guides that a general web search can surface,"** not as a designed cross-section of the roughly 198 ABA-accredited law schools. [4] I want to be direct that no claim about "typical" or "common" law school practice can be supported by a sample assembled this way, and the Deep Analysis in the prior turn was deliberately scoped to avoid such claims for that reason.

(2) Definition and Placement of "Borderline Cases"

I used "borderline" imprecisely in the prior response and should define it now. A case is **borderline** when the retrieved primary text does not use the same categorical language as the coding scheme, such that assigning it to a coding cell requires an inference about what the drafters most likely meant, rather than a direct transcription of the drafters' own stated category.

This is distinct from a **gap** (no text at all) and distinct from a **clear case** (the text uses language that maps onto the coding variable without interpretation, e.g., Columbia's explicit "Default Prohibition" heading). [\[5\]](#)

Three borderline determinations from the prior inventory, made explicit here:

- **Fordham and Mitchell Hamline, coded as "governed through a plagiarism definition rather than a standalone AI framework."** Fordham's Standards of Academic Integrity states students must complete work "without assistance from external resources (including any artificial intelligence or machine-learning tools) unless directly authorized by the instructor," embedded within its general plagiarism/cheating definitions rather than in a separate AI-specific section. Mitchell Hamline's Student Code of Conduct defines plagiarism to include use of "an artificial intelligence content-generator's words or ideas" without documentation, again inside the general plagiarism definition rather than a freestanding AI policy. [\[6\]](#) [\[7\]](#) [\[8\]](#) The inference required was to characterize this drafting choice (AI folded into an existing doctrinal category) as functionally different from schools that created a new, separate instrument (Berkeley, Columbia). That characterization is reasonable but is my synthesis, not the schools' own self-description — neither institution's text says "we have chosen to regulate AI via the plagiarism doctrine rather than through a freestanding policy."
- **American University Washington College of Law, coded as "full delegation to instructor discretion."** The only retrieved text is a secondary description stating the Academic Integrity Code "is not prohibited" for student AI use and there is "no one-size-fits-all" rule. This is itself a paraphrase from a library guide, not a quotation of AU's own policy instrument, so the "full delegation" classification is an inference drawn from a secondary source's summary — one further step removed from the primary text than the other codings, and accordingly a weaker-confidence classification than, for example, Berkeley's or Columbia's. [\[9\]](#) [\[10\]](#)
- **CTLS's institutional identity relative to Georgetown Law.** In the original inventory I flagged, but did not resolve, whether the "Center for Transnational Legal Studies" policy should be understood as Georgetown University Law Center's home-campus JD policy. CTLS is a consortium program historically administered in partnership with Georgetown Law but serving students from multiple member law schools; nothing in the retrieved text confirms it governs Georgetown's standard JD coursework. I flagged this rather than resolving it by assumption, and the same caution should extend to any comparative statement that treats "Georgetown" as a coded data point — the coded row should be understood as "CTLS," not "Georgetown Law" generally.

Where a determination was **not** borderline in this sense — for example, coding Columbia's exam rule as an absolute default prohibition — I relied on direct quotation ("Use of Generative AI is prohibited in (a) any exam or final paper") requiring no inference about drafters' intent. [\[11\]](#)

(3) Mapping the Compliance Architectures to Established Regulatory

Typologies

Legal and regulatory theory offers several established typologies that can be applied to the coded dataset without importing a restrictiveness ranking. Two are especially well-suited here because they are used to classify *how* a rule is structured, not *how strict* it is.

A. Rules versus standards. This is a foundational distinction in regulatory design theory: a *rule* specifies ex ante, with relative determinacy, exactly what conduct is required or forbidden, minimizing interpretive cost at the point of application; a *standard* specifies a more general norm to be applied and interpreted ex post, trading determinacy for flexibility. [12] [13] Applying this distinction to the coded schools:

- **Rule-like drafting:** Berkeley's policy enumerates specific prohibited activities in bulleted form — "Asking an AI tool to brainstorm a paper topic or thesis (prohibited conceptualizing)," "Asking an AI tool to propose an organizational structure for a paper (prohibited outlining)" — which is closer to what the Legal Theory Lexicon separately identifies as a *catalog* (an enumerated, illustrative list of covered instances) layered onto a rule. [14] Columbia's prohibition is similarly rule-like: it specifies the covered conduct ("any exam or final paper," "aid in drafting any part of work submitted for credit") with minimal need for ex post interpretation. [15]
- **Standard-like drafting:** Chicago's policy asks whether AI use "would constitute academic plagiarism if the generative AI were a human author whose work was used without attribution" — this requires an evaluator to apply an existing, more general doctrine (plagiarism) to a new fact pattern, which is the defining feature of a standard rather than a rule. [16]
- **Hybrid drafting:** Stanford's policy combines a standard-like default ("support their learning and to aid in the development or refinement of their own ideas") with a rule-like bright-line trigger (drafting or exam use requires advance written disclosure and authorization, full stop). [17]

This mapping is useful because it reclassifies the schools along an axis — interpretive determinacy — that is orthogonal to substantive restrictiveness. Chicago and Berkeley reach similarly restrictive outcomes for exams, but Berkeley does so via rule-drafting and Chicago via standard-drafting, which has different downstream implications for predictability and adjudication cost that are independent of how permissive or restrictive either school is.

B. Default versus mandatory rules. A second typology from private-law and regulatory-design scholarship distinguishes *default rules* (which apply absent a contrary specification by the relevant party and can be "contracted around") from *mandatory rules* (which bind regardless of the parties' preferences). [18] This typology maps unusually well onto the dataset because nearly every coded school's policy operates on **two different axes simultaneously**, with different default/mandatory characterizations on each:

- On the **institution-to-instructor axis**, the core prohibitions in Berkeley, Columbia, Chicago, and USD's policies function as *default rules*: they apply "in the absence of an explicit policy for the class set by the instructor" and instructors may "deviate from the default rule, provided that they do so in writing and with appropriate notice". [19] This is the textbook structure of a contractible default.
- On the **instructor-to-student axis**, the same rule typically functions as a *mandatory rule*: an individual student cannot waive it by personal disclosure or good-faith intent. Columbia states this explicitly — use is prohibited "even if the use is fully documented" — meaning student-side disclosure does not achieve the kind of "contracting around" that an instructor's written syllabus notice can achieve. [15] The same rule is therefore simultaneously a default

(relative to the instructor, who holds override authority) and mandatory (relative to the student, who does not).

This two-axis default/mandatory structure also reveals a further typological distinction worth naming: **baseline polarity**. Some schools set a *restrictive default that instructors may loosen* — Berkeley, Columbia, USD, and (for exams) Chicago all begin from prohibition and require an affirmative, written instructor act to permit more. [20] [21] [22] CTLS inverts this baseline: its default is permissive ("Unless explicitly prohibited... students may reasonably assume AI tool usage is permissible... provided proper attribution is maintained"), requiring an affirmative instructor act to restrict. This is precisely the "sticky default rule" phenomenon described in contract-default-rule scholarship: because opting out requires effort (drafting and publishing written notice), the choice of baseline polarity itself functions as a substantive lever even holding the range of permissible instructor choices constant. [23] A law school that sets a restrictive default will likely see more restrictive *outcomes* in courses where an instructor never gets around to writing an alternative syllabus policy than a law school with an identical range of permissible instructor choices but a permissive default — independent of any school's stated substantive preference.

C. Delegation/discretion typology. A third, cruder but still useful axis distinguishes fully centralized rules (no delegation), bounded delegation (default-with-override, as above), and full delegation (no institutional default at all, purely instructor-determined). Harvard, Yale, NYU, and American University's general statements fall into the third category on the retrieved evidence — there is no institutional default to override because the institution has not stated one at the law-school level, only a general instruction that instructors must state their own policy. [9] [24] [25] [26] Fordham and Mitchell Hamline sit in an intermediate position: a mandatory institutional floor (the general plagiarism rule) exists, but within that floor, an instructor-permission trigger reintroduces a default-like feature at the point where a specific use is or is not "authorized". [7]

(4) Penn State's Law Schools

Penn State operates two connected law school locations under Penn State Dickinson Law (Carlisle and University Park, the latter having previously operated as a separate "Penn State Law"). [27] Research for this response did not identify a Penn State Dickinson Law-specific coursework or exam AI policy instrument comparable to Berkeley's, Columbia's, or Stanford's. What was identified instead:

- **A University-wide (not law-specific) academic integrity policy on AI**, published by Penn State's central Office of Academic Integrity, which states directly: "Students may not use generative AI tools to complete multiple-choice, matching, fill-in the blank, open-ended, or essay exam questions". [28] This is a comparatively rule-like, bright-line exam prohibition — more specific in its enumeration of exam question types than the general university-level statements found for Harvard, Yale, or NYU. The same page instructs students to "read the course syllabus and assessment instructions," since "in some classes, students may use AI tools. In some classes, students may not," and directs disclosure practice where AI is permitted: naming the tool used, citing AI-derived content in-text, and listing AI tools as a "Sourced Tool" at the end of the paper. [29] [30] [31] This places Penn State's baseline

architecture in the same "default rule with instructor override" family described above, but notably with a *university-level mandatory floor specifically for exams* that is more explicit and more granular than the comparable university-level statements at Harvard, Yale, or NYU.

- **A Dickinson Law Library research guide** ("Generative AI for Legal Educators"), which is explicitly an informational resource for faculty rather than a binding rule, and which links out to the university's general AI Hub guidelines, sample syllabus statement templates, and training resources rather than stating a law-school-specific rule of its own. [32] [33] This guide's existence indicates active institutional engagement with AI pedagogy at the law school but does not itself constitute retrievable law-school-specific policy text, and I am not treating it as such.
- **An institutional practice signal worth noting separately from any conduct rule:** in February 2026, Penn State Dickinson Law's Montague Law Library announced a partnership with Harvey (a legal-domain generative AI platform) under Harvey's Law School Partnership Program, to be deployed "in classes, experiential learning, and clinics". [34] [35] The library's own announcement includes a caution rather than a permissive framing: "Generative AI carries a risk of hallucinations, even in a closed system. Always verify your results," and notes that "Penn State policy prohibits uploading client-sensitive or private data to Harvey". [36] This is analogous in kind (though not necessarily in scale) to UT Austin's licensing arrangement discussed previously — an infrastructure-provisioning decision that sits alongside, but is not itself, a coursework conduct rule. [37]

On the retrieved evidence, Penn State Dickinson Law is best classified in the same structural family as Harvard, Yale, NYU, and American University: **a general (here, unusually specific-for-exams) university-level default, with law-school-specific implementation left to individual course syllabi**, rather than a standalone law-school AI policy instrument of the kind Berkeley, Columbia, Chicago, Stanford, USD, or CTLS have issued. It differs from those other four schools in that the university-level exam rule is drafted with rule-like specificity (naming exam question formats) rather than a general instruction to "adhere to academic integrity policies."

(5) Faculty Autonomy in Coursework, Exam, and Paper AI Policy

Faculty autonomy is not a peripheral feature of these policies; on the coded evidence, it is close to the central organizing design choice, and it operates differently depending on where a given school's baseline sits.

Autonomy as bidirectional override. Several schools' commentary explicitly authorizes instructors to move the applicable rule in *either* direction from the institutional default. Columbia's commentary states instructors "may... permit specific uses of Generative AI if students show how and where they rely on it; or alternatively they may prohibit uses otherwise permitted under this policy" — a genuinely two-way discretion. [38] Berkeley similarly allows instructors to "deviate from the default rule for courses designed intentionally to teach AI fluency (or for other courses for which the instructor decides a distinct rule is pedagogically appropriate)", without textually limiting deviation to a single direction. [22] Chicago's exam prohibition applies "unless the instructor specifies otherwise", again framed as open in direction.

Autonomy as one-directional loosening only. Other schools' textual structure permits only a loosening of a restrictive default, not further tightening — because the default is already the maximally restrictive institutional floor. USD's rule states "no AGC may be included in written work submitted for credit... unless the instructor explicitly permits such use in writing": the instructor's only textually contemplated act is to grant permission, not to impose additional restriction beyond the default prohibition (though nothing in the text forecloses that either). [39] Fordham's structure is comparable: the university default prohibits unauthorized use "unless directly authorized by the instructor", with authorization being the only instructor act named. [7]

Autonomy as one-directional tightening only. CTLS presents the inverse: because the institutional default is permissive, the only instructor act contemplated by the text is restriction — "If a faculty member decides to limit or prohibit the use of AI, they shall provide clear instructions on what is permitted or prohibited". [40] Here, faculty autonomy is exercised by narrowing an otherwise-permissive baseline, the reverse operation from USD or Fordham.

Bounded autonomy — floors that instructor discretion cannot cross. In every school where the text addresses this, instructor discretion is capped by a background norm the instructor cannot waive on a student's behalf. Stanford states that "in all cases, the student's use of AI cannot contravene standard academic norms concerning plagiarism and accuracy, and in coursework involving the practice of law, AI must meet all applicable rules of professional responsibility" — meaning an instructor's permissive course policy is itself bounded by plagiarism doctrine and (in clinical settings) professional-responsibility rules that no instructor can contract around. [41] Chicago's academic-plagiarism-equivalence test operates the same way: whatever an instructor permits, "the default policy sweeps more broadly than a prohibition against academic plagiarism," and "a student cannot avoid limits on the use of generative AI that would otherwise apply by attributing their work to generative AI". [42] [43] This indicates that faculty autonomy, across every school where it is textually elaborated, is designed as *bounded* discretion — a zone of instructor choice nested inside an outer, non-waivable institutional or professional floor — rather than unlimited course-by-course sovereignty.

A further, previously unremarked layer: sub-law-school program-level autonomy. The evidence also shows autonomy operating below the level of the individual instructor, at the level of a clinical program. Stanford Law's general default policy applies "to all Stanford Law School courses, including clinics", but a specific clinic within the law school (the Juelsgaard Intellectual Property and Innovation Clinic) has published its own, more restrictive syllabus language stating that "as a default, you should not use ChatGPT or any other generative AI tools in completing any of your clinic assignments," citing sanctions imposed on attorneys who used generative AI to prepare briefs. [44] [45] This indicates a four-tier discretion hierarchy at least at Stanford — university guidance, law-school default, clinical-program-specific policy, and (within that clinic's own language) further individual-supervisor discretion — rather than the two-tier (institution/instructor) structure that the coding scheme's "rule-setting architecture" variable, as originally defined, was built to capture. This is a limitation of the original coding scheme worth flagging directly: it did not have a category for program-level (as distinct from course-level) policy variation, and the Stanford clinic example shows that such variation exists in at least one school in the sample.

Why the direction of the "no default" cases (Harvard, Yale, NYU, American University, Penn State) matters analytically. In these five schools, the retrieved evidence shows no institutional substantive default at the law-school level at all — only a procedural requirement

that *some* instructor-level policy be stated in writing, plus (at Penn State) a university-wide exam-specific floor. [\[28\]](#) [\[46\]](#) This is a qualitatively different design choice from "default-with-bidirectional-override": it is closer to pure delegation, where the only thing mandated centrally is the *existence* of a stated policy, not its content (subject to Penn State's exam floor as a partial exception). Yale's framing captures this directly: "each course at Yale sets its own policies for using AI. Using AI when it's not authorized in a course constitutes academic dishonesty" — the substantive content of "authorized" is entirely instructor-determined, with the institution's role limited to enforcing whatever the instructor decided, after the fact, as a matter of general academic-integrity doctrine. [\[25\]](#)

References

1. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
2. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
3. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
4. [ABA-Approved Law Schools](#)
5. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
6. [Standards of Academic Integrity - Fordham University](#)
7. [Standards of Academic Integrity - Fordham University](#)
8. [Standards of Academic Integrity - Fordham University](#)
9. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
10. [WCL's Policies - Guidelines for Law Journals and Student Authors: Plagiarism and AI Use - LibGuides at American University Washington College of Law](#)
11. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
12. [\[PDF\] Rules vs. Standards in Private Ordering](#)
13. [\[PDF\] Rules vs. Standards in Private Ordering](#)
14. [Artificial Intelligence Policy - UC Berkeley Law](#)
15. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
16. [4.9 Law School Policy on Generative AI](#)
17. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
18. [A Theory of Mandatory Rules: Typology, Policy, and Design | Texas Law Review](#)
19. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
20. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
21. [Artificial Intelligence Policy - UC Berkeley Law](#)
22. [Artificial Intelligence Policy - UC Berkeley Law](#)
23. ["On the Stickiness of Default Rules" by Omri Ben-Shahar ...](#)
24. [Frequently Asked Questions About Teaching and AI](#)
25. [AI Guidance for Students | Poorvu Center for Teaching and Learning](#)
26. [Guidelines for Using ChatGPT and other Generative AI tools at Harvard | Office of the Provost](#)
27. [Penn State Dickinson Law](#)
28. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
29. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
30. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
31. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)
32. [Getting Started - Generative AI for Legal Educators - Dickinson Law Library at Penn State Dickinson Law](#)
33. [Penn State Resources - Generative AI for Legal Educators - Dickinson Law Library at Penn State Dickinson Law](#)
34. [Penn State Dickinson Law joins Harvey's Law School ...](#)
35. [Penn State Dickinson Law joins Harvey's Law School ...](#)
36. [Penn State Dickinson Law joins Harvey's Law School ...](#)

37. [University of Texas Law Alters AI Policy to Stop Skills Loss \(3\)](#)
38. [Columbia Law School Interim Policy on Generative AI | Columbia Law School](#)
39. [AI Class Policies - Generative AI Tools for USD Law Faculty - Guides at University of San Diego Legal Research Center](#)
40. [Policy on the Use of Artificial Intelligence in Academic Work](#)
41. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
42. [4.9 Law School Policy on Generative AI](#)
43. [4.9 Law School Policy on Generative AI](#)
44. [Use of Generative AI Technology - Student Affairs - Stanford Law School](#)
45. [Clinic, Law School, and University AI Policies and Syllabus Language - Juelsgaard Intellectual Property and Innovation Clinic - Stanford Law School](#)
46. [Artificial Intelligence - Academic Integrity \(psu.edu\)](#)