

Law 580: Legal Writing in the AI Era

Thursdays at 9:40 a.m. – 11:40 a.m. (August 20 – October 1)
Antonin Scalia Law School, George Mason University, Fall 2026

Instructor: Michelle Trumbo
Interim Director, Library & Technology
Assistant Professor and Assistant Director, LRWA Program

Office Hours: Room 243 (in the Law Library's Administration suite)
Thursdays, 12 p.m. – 2 p.m. or by appointment

Contact Info: mtrumbo2@gmu.edu | (703) 993-8103

Course Description

This upper-level course explores the role of artificial intelligence in legal writing and law practice. Students learn to evaluate and critically engage with AI tools that generate, edit, and analyze text, while strengthening their own professional writing, analytical reasoning, and judgment. The course emphasizes both the opportunities and the limitations of AI, including issues of accuracy, bias, and the lawyer's enduring responsibility for their work product.

Where AI is concerned, the central skill is supervision, and supervision runs in both directions. Students learn to command the *inputs* — choosing the right tool for the task and instructing it with precision — and to discipline the *outputs* — verifying, correcting, and revising what the machine returns. A lawyer who cannot do the first commissions bad work; a lawyer who cannot do the second signs it.

Through guided exercises and a sustained drafting sequence built around a single hypothetical, students integrate AI-assisted work into multiple forms of legal writing, from research planning through persuasive and transactional drafting. Special attention is paid to ethical and professional responsibility, including the duties of competence, confidentiality, candor, supervision, and client communication, and to the rapidly developing case law on sanctions for misuse of AI in court filings.

AI Policy

Students are expected to use AI tools throughout this course, and to disclose that use: every submitted exercise and portfolio entry identifies which tools were used, for what, with the relevant prompts logged. Only tools approved for coursework may be used with assignment materials. No real client information, or any confidential information from employment or clinical work, may be entered into any AI tool for any purpose in this course. Undisclosed AI use is treated as a violation of the Academic Regulations, exactly the way undisclosed AI use in practice is treated as a violation of the duty of candor.

Course Objectives

By the end of the course, students will be able to:

- 01.** explain how generative language models produce text and anticipate their characteristic failure modes;
- 02.** select the appropriate AI tool for a given task using a defensible framework (e.g., task fit, grounding, confidentiality, cost, and verifiability) and “bench test” an unfamiliar tool before relying on it;
- 03.** construct, iterate, and document effective prompts for legal research, drafting, and editing, including grounded (“closed universe”) prompting and adversarial prompting;
- 04.** judge when AI can responsibly be used in legal writing and when reliance would be inappropriate;
- 05.** apply a systematic verification protocol to every AI output, covering citations, holdings, jurisdiction, scope, and currency;
- 06.** detect and correct hallucinated authority, subtle substantive errors, and embedded bias in AI outputs;
- 07.** revise AI-assisted drafts to professional standards across objective, persuasive, client-facing, and transactional genres; and
- 08.** practice within the ethical duties governing AI use (i.e., competence, confidentiality, candor, supervision, fees, and client disclosure) and develop a professional identity as an AI-mediated legal writer.

Curriculum Map

Each objective is introduced early, reinforced across the middle weeks, and applied and assessed in the capstone. Reading down a column shows what a given week contributes; reading across a row shows how a skill matures.

Course objective	W1	W2	W3	W4	W5	W6	W7
01. Explain how generative models work & anticipate failure modes	I	D		D		D	A
02. Select the right tool with a defensible framework; “bench test” new tools	I		D	D			A
03. Construct, iterate & document effective prompts (grounded, adversarial)		I		D	D	D	A
04. Judge when AI use is responsible and when reliance is inappropriate	I	D	D	D	D	D	A
05. Apply a systematic verification protocol to every AI output		I	D	D	D	D	A
06. Detect & correct hallucination, subtle error, and embedded bias		I		D	D	D	A

Course objective	W1	W2	W3	W4	W5	W6	W7
07. Revise AI-assisted drafts to professional standards	I				D	D	A
08. Practice within ethical duties & articulate a professional identity	I		D		D		A
The week's focus:	<i>Know the machine & market</i>	<i>Craft inputs; doubt outputs</i>	<i>Know the rules</i>	<i>Plan with AI</i>	<i>Draft & advise</i>	<i>Persuade & stress-test</i>	<i>Tying it all together</i>

Key: **I** introduced **D** developed **A** applied & assessed

Credit and structure

1 credit delivered over seven weeks (August 20, 2026 – October 1, 2026). Class will meet for 1 two-hour session per week, with outside work including readings, exercises, and portfolio entries.

Prerequisites

Legal Research, Writing, and Analysis (LRWA I and II) or equivalent foundational legal writing courses. No prior experience with generative AI is assumed; students who have used these tools extensively will find a structured framework for what they may already do informally.

Readings

Readings vary by weekly topic; a full annotated bibliography will be posted on Canvas. Additionally, assigned reading will periodically include two or three current entries from a public AI-sanctions tracker.

Assessment

Grades are weighted as follows:

- **Weekly exercises (20%).** Short, structured drafting or verification tasks tied to each week's topic; graded low pass / pass / high pass.
- **Class participation (10%).**
- **Process portfolio (45%).** A collection of annotated drafts, prompt logs, revision histories, tool-choice justifications, reflections, and decision rationales generated throughout the course.
- **Final integrated sequence document (20%).** The polished memo, brief excerpt, and transactional provision from the integrated hypothetical.
- **Prompt playbook (5%).** A curated personal library of tested, annotated prompts, begun in Week 2 and submitted with the portfolio.

Skills That Build Across the Course

These five competencies are not confined to a single week. Instead, they develop as the course progresses, and each weekly module closes with a short block of recurring skill work.

- **Verification.** Introduced in Week 2; drilled every subsequent week. By Week 7 it should feel automatic.
- **Prompt craft.** Introduced in Week 2; sharpened in a short prompt clinic at the top of each of Weeks 4–7, each aimed at that week’s theme; consolidated in the prompt playbook.
- **Bias.** Introduced in Week 2; treated substantively in Week 5 (framing of facts and audience), Week 6 (rhetorical assumptions in persuasion), and Week 7 (boilerplate that bakes in assumptions).
- **Confidentiality and data security.** Introduced with the tool-selection framework in Week 1; treated fully in Week 3; revisited in Week 7, where deal documents raise distinct disclosure concerns.
- **Professional identity.** A Week 1 position statement on what AI should and should not do in legal writing is returned in Week 7 and revised considering the course’s content.

Weekly Modules

Week 1: How the Machines Work and Which Tool to Use

The course opens with the mechanics: how large language models work, why they generate coherent but sometimes unreliable text, and what that means for legal writers. Discussion considers the potential benefits of AI in practice (drafting efficiency, brainstorming, stylistic range) alongside the most common pitfalls (fabrication, superficial analysis, overreliance).

The next part of the session is devoted to tool selection. This will be taught as a decision framework rather than a product tour, because the products will change and the framework will not. Students learn the tool categories (general-purpose frontier models; retrieval-grounded legal research tools; drafting-focused tools; litigation-support and agentic tools) and a four-question test applied before any task:

- **Retrieval or generation?** Does this task require finding real authority, or producing and shaping text? Never ask an ungrounded general-purpose model for authority.
- **Closed or open universe?** Will the tool work from documents I supply (far more reliable), or from the world at large?
- **What confidentiality tier is the data?** Consumer tools, enterprise tenancy, and firm-approved systems sit at very different risk levels. This will be developed further in Week 3.
- **How will I verify the output and what will verification cost?** If verification costs more than doing the work directly, the tool is the wrong choice.

Because new tools arrive constantly, students also learn a “bench test” protocol: a personal suite of known-answer tasks (e.g., a citation you can check, a doctrine you know cold, a document you have already analyzed) run against any new tool before trusting it with real work.

Week 2: The Craft of the Prompt and the Discipline of Verification

This week pairs the two halves of supervision: instructing the machine well and refusing to trust what it returns. They are taught together because they are the same cognitive work, simply pointed in opposite directions.

The first part of class is a working session on prompt construction for legal tasks:

- **Anatomy of an effective legal prompt.** Role, context, task, constraints, and format (with worked examples showing how each added element changes the output). The controlling analogy: a prompt is an assignment memo to a junior associate, and vague prompts/assignments produce bad work product in both cases.
- **Grounded (“closed universe”) prompting.** Supplying the source documents and confining the model to them is the single most effective hallucination-reduction technique available to a practicing lawyer. Also employ a quote-and-pinpoint approach: instruct the model to quote its sources verbatim with pinpoint citations, because what it cannot quote it may be inventing.
- **Iteration and decomposition (“chunking”).** Refining a prompt across several passes; breaking a large task into a prompt chain (issues → outline → section drafts → edit passes) rather than asking for a finished product in one shot.
- **Eliciting doubt.** Instructing the model to flag uncertainty, state assumptions, and argue against its own answer. This will be an introduction to the adversarial prompting techniques further developed in Week 6.
- **Multi-draft comparison.** Generating several drafts and comparing them. The takeaway here is that generation is cheap, fast, and that AI output is a draft, not the draft.
- **What prompting cannot fix.** No prompt makes an ungrounded model a reliable source of authority. Prompt craft narrows the error rate; it never eliminates the duty to verify.

Later in this class we will turn to reliability. Students examine high-profile sanctions cases, study why models fabricate plausible but nonexistent authority, and work through AI-produced briefs riddled with errors. The class then builds a broader typology of AI errors (e.g., misinterpreting the question, drifting outside the assigned scope, citing real but overruled or distinguished authority, applying the wrong jurisdiction’s rule, subtle anachronism, etc.). From it, each student builds the accuracy-check protocol they will reuse every subsequent week. A short closing segment introduces bias in training data and outputs a topic that gets greater treatment in Weeks 5–7.

Week 3: Ethics, Confidentiality, and Professional Responsibility

This week explores the professional and ethical landscape of AI in legal practice: the ABA Model Rules as applied to AI use, ABA Formal Opinion 512, and the growing body of state bar opinions. Topics include the duty of competence (Rule 1.1), candor toward the tribunal (Rule 3.3), supervision of nonlawyer assistance including technology (Rules 5.1 and 5.3), fees and the billing of AI-assisted time (Rule 1.5, as addressed in Opinion 512), and the developing landscape of court-by-court AI rules, including the pre-filing practice of checking the assigned judge’s standing order on AI use, which now varies dramatically across courtrooms.

A portion of the class addresses confidentiality and data security. That is, what may be entered into which kind of tool, the difference between consumer and enterprise-tenancy products, redaction and

anonymization habits, and the practical decision tree students will face on their first day of practice. A closely related sub-segment addresses privilege and work product: whether feeding client documents to a hosted tool waives attorney-client privilege or work-product protection, the developing case law, and the safeguards firms are adopting (no-training contractual terms, anonymization protocols, approved-tool lists). A further sub-segment addresses disclosure to clients and informed consent (NB: this is distinct from disclosure to tribunals) under Florida Opinion 24-1 and similar state guidance, and the interaction with Rule 1.4.

A final segment addresses AI as adversary. This is the attorney's duty to catch the other side's AI errors, not just one's own, as courts begin to sanction lawyers who fail to detect fabricated authority in opposing briefs and the rapidly growing volume of AI-assisted pro se filings, with its implications for opposing counsel and for the courts.

Week 4: AI in Research and Planning (Integrated Sequence I)

During Weeks 4 through 6 students will work a single hypothetical from research plan to objective memo to persuasive brief, and the deliverables become the body of the process portfolio. The session addresses where in the writing process AI supports research and planning (e.g., idea generation, narrowing broad questions, flagging gaps, generating skeletal outlines) and then confronts AI's limits as a research tool, anchored empirically in the Stanford RegLab findings on hallucination rates in the leading legal research platforms. A short segment covers agentic and deep-research tools (i.e., systems that plan and execute multi-step research and drafting on their own). The key takeaway is the longer the leash, the more the errors compound, and the duties of verification and supervision do not scale down as autonomy scales up.

Week 5: Objective Drafting and Client Communication (Integrated Sequence II)

The week shifts from planning to drafting the internal memorandum. Students examine how AI assists with neutral language and structural experimentation, while learning its habits: overgeneralization, inconsistent organization, missed nuance in rule application. Editing and polishing follow — where AI functions well as an editorial assistant — with students comparing the AI's edits against their own and deciding, edit by edit, which changes sharpen the analysis and which flatten it. This week makes explicit the critique-over-composition frame: the lawyer's distinctive contribution is evaluating prose, not generating it.

This class will also extend objective writing to client communication, the most common writing a junior associate produces and a type of legal writing too often treated as an afterthought. Students examine AI's use in client emails, status updates, demand letters, and explanations for non-lawyer audiences and its distinctive risks: projecting false confidence, softening bad news into misleading reassurance, eliding the qualifications a lawyer would include. A short segment addresses citation and attribution norms in the AI era. Exploring whether AI-assisted research should be cited, whether AI-generated language should be flagged, and where a writer crosses from "used a tool" to "used a ghostwriter." In some situations, these questions are unsettled and the choices students make will signal their individual professional norms/identity.

Week 6: Persuasive Writing, Sycophancy, and Adversarial Prompting (Integrated Sequence III)

The course turns to persuasive writing, where the week's central problem is **sycophancy**: AI tools are trained to be agreeable, so they will almost always tell a writer her argument is strong, her framing persuasive, her conclusion sound. For a lawyer building a brief, that is precisely the wrong feedback. The class works through the research on sycophancy and its implications, then examines how AI can shift tone, structure, and emphasis when a document's purpose changes from internal advice to external advocacy. We will also discuss where AI's sense of audience, rhetorical strategy, and persuasive nuance falls short.

Week 7: Transactional Drafting, the Baseline Revisited, and Synthesis

The first part of class addresses transactional drafting which is distinct enough from litigation to warrant its own treatment. Students explore how AI assists with contracts, corporate policies, and compliance documents, where precision, structure, and anticipation of contingencies are paramount: its strengths (first drafts, boilerplate suggestion, inconsistency detection) and its limits (no business context, no feel for deal dynamics, and a well-documented tendency to import unenforceable or jurisdictionally inappropriate terms). Confidentiality issues return in transactional form through deal documents and data that raise disclosure concerns beyond a litigation context.