

Zen Legal: AI-Assisted Legal Research and Analysis

Technical Report v2025.10

Zach Kelling
Zen LM Research Team
research@zenlm.org

October 2025

Abstract

We present Zen Legal, a domain-adapted language model for legal research, document analysis, and contract review. Zen Legal is *not* a from-scratch model: it is a supervised fine-tune of the openly licensed **Qwen3-8B** base model (developed by Alibaba and released under the Apache-2.0 license), adapted to the legal domain using the Zen-Pro fine-tuning recipe. Building on Qwen3-8B’s general reasoning and instruction-following capabilities, Zen Legal targets the fundamental requirements of legal AI: citation-grounded generation, multi-jurisdiction reasoning, precise language preservation, and privilege protection. This report describes the legal domain adaptation methodology, the citation-grounding and retrieval architecture, jurisdiction-aware reasoning, and professional responsibility guardrails. We deliberately do not report headline accuracy figures on legal benchmarks here; rather, we describe the adaptation qualitatively and point to the public benchmarks (LegalBench, CUAD, and others) against which legal language models should be evaluated under transparent, reproducible protocols.

1 Introduction

Legal practice demands precision at a level that distinguishes it from most AI application domains. A misquoted statute, a missed precedent, or a hallucinated citation can undermine a legal argument, expose a client to liability, or trigger bar discipline for the attorney who relied on it.

Zen Legal is built on the premise that the primary role of AI in legal contexts is *augmentation*, not replacement—reducing research time from days to hours, identifying overlooked precedents, flagging inconsistencies in contract drafts, and generating first drafts that attorneys review and refine. This framing shapes every design decision:

1. **Citation-grounded generation:** Every legal claim cites a specific case, statute, or regulation. Uncited assertions are prohibited.
2. **Jurisdictional awareness:** Legal rules vary by jurisdiction; Zen Legal explicitly identifies which jurisdiction’s law it is applying and flags conflicts.
3. **Uncertainty disclosure:** The model expresses calibrated uncertainty about unsettled legal questions rather than stating contested doctrine as established law.
4. **Privilege protection:** Attorney-client privileged content is handled with enhanced access controls and never leaves the client environment.
5. **Professional responsibility compliance:** Outputs comply with ABA Model Rules and applicable state professional conduct rules.

2 Legal Domain Adaptation

2.1 Base Model and Licensing

Zen Legal is a domain fine-tune, not a model trained from scratch. The base model is **Qwen3-8B** [1], an 8-billion-parameter dense transformer developed and released by Alibaba under the Apache-2.0 license. We chose Qwen3-8B for its strong general reasoning, long-context handling, and instruction-following behavior, together with a permissive license that allows redistribution of derived weights. All capability claims in this report should be read as the result of *adapting* this base model to the legal domain; the underlying architecture, pre-training, and general-domain knowledge are inherited from Qwen3-8B and are the work of its original authors, whom we credit accordingly. Zen Legal is *not* based on a bespoke “Zen MoDE” or mixture-of-experts backbone.

2.2 Fine-Tuning Corpus

Zen Legal is fine-tuned on a large legal corpus assembled from the following sources. The corpus is used for continued pre-training and supervised fine-tuning on top of Qwen3-8B; the relative weights below describe sampling proportions during adaptation, not pre-training-scale token budgets.

Table 1: Legal fine-tuning corpus composition (relative sampling share during domain adaptation of Qwen3-8B)

Source	Coverage
Federal case law (PACER, CourtListener)	1789–2025
State case law (all 50 states)	Selective post-1970
Federal statutory code (USC)	Current
Code of Federal Regulations (CFR)	Current
Law review articles (HeinOnline)	1820–2025
International law (UN, EU, OECD)	1945–2025
Legal contracts and agreements	Commercial
Legal treatises and study materials	Current
Synthetic legal QA pairs	All jurisdictions

2.3 Citation Grounding Architecture

Legal citations require hyper-precise attribution: case name, reporter, volume, page, pinpoint page, court, and year. Zen Legal’s citation generation combines:

1. **Retrieval:** A dense retrieval index over 48 million legal documents retrieves candidate authorities.
2. **Verification:** Retrieved documents are verified against a citation registry that checks existence and accuracy of pinpoint cites.
3. **Shepardization:** Citations are checked against a legal validation service to flag overruled, distinguished, or limited precedents.
4. **Attribution:** The model generates text with inline citation markers that link to verified source documents.

Citation accuracy—whether a generated citation exists and supports the stated proposition—is the central evaluation target for this pipeline and should be measured on a held-out set

reviewed by licensed attorneys. We report no headline accuracy figure here; the retrieval-plus-verification design is intended to make uncited or unverifiable assertions structurally impossible rather than to maximize a single benchmark number.

3 Multi-Jurisdiction Reasoning

3.1 Jurisdiction Classification

Every legal query is classified by the applicable jurisdiction(s) before substantive analysis:

- Primary jurisdiction: The jurisdiction whose law governs the issue.
- Secondary jurisdictions: Related jurisdictions consulted for persuasive authority.
- Conflict detection: When jurisdictions differ, the model explicitly identifies the split.

$$P(j \mid q) = \text{softmax}(\mathbf{W}_j \cdot \phi(q)) \quad (1)$$

where $\phi(q)$ is the contextual representation produced by the fine-tuned Qwen3-8B backbone and $\mathbf{W}_j$ is a lightweight jurisdiction-classification head learned during adaptation.

3.2 Circuit Splits and Unsettled Law

When a question implicates a circuit split or unsettled area of law, Zen Legal:

1. Identifies the split (e.g., “The Third and Ninth Circuits have held X, while the Fifth Circuit holds Y”).
2. States which rule applies in the user’s jurisdiction.
3. Notes any Supreme Court certiorari activity or pending resolution.
4. Advises the user to verify current status given the dynamic nature of the area.

4 Contract Analysis

4.1 CUAD Clause Extraction

The Contract Understanding Atticus Dataset (CUAD) [3] is the standard public benchmark for commercial contract clause extraction. It annotates 41 clause types across hundreds of commercial agreements, including governing law, limitation of liability, indemnification, non-compete, termination for cause, assignment, IP ownership, audit rights, change of control, warranty duration, renewal term, most-favored-nation, exclusivity, anti-assignment, and force majeure clauses. Zen Legal’s domain adaptation explicitly targets this clause taxonomy: the fine-tuning corpus includes contract text labeled with these clause categories, and the model is trained to extract and span-tag clauses by type.

We report no per-clause accuracy figures in this report. CUAD is a public benchmark with a published evaluation protocol, and we recommend that any deployment of Zen Legal be evaluated against it directly, with the evaluation harness, model version, and prompt configuration disclosed so that results are reproducible and comparable to the published CUAD baselines.

4.2 Contract Risk Scoring

Beyond clause extraction, Zen Legal generates risk assessments for each contract:

- **Clause-level risk:** Each extracted clause is scored on a 1–5 risk scale relative to market standard terms.
- **Missing clause detection:** Identifies standard clauses absent from the agreement (e.g., no limitation of liability, no governing law clause).
- **Inconsistency detection:** Flags internally inconsistent provisions (e.g., conflicting notice periods in different sections).
- **Redline suggestions:** Proposes alternative language for high-risk clauses with explanation.

5 Evaluation Methodology

Rather than report headline scores, this section names the public benchmarks against which a legal language model such as Zen Legal should be evaluated, and the protocol we recommend. Because Zen Legal is a fine-tune of Qwen3-8B, the appropriate comparison is between the adapted model and the unmodified Qwen3-8B base, holding the evaluation harness and prompts fixed, so that any measured difference is attributable to the legal adaptation rather than to a different backbone.

5.1 Recommended Public Benchmarks

- **LegalBench** [2]: a collaboratively built benchmark of 162 legal-reasoning tasks spanning contract, constitutional, administrative, criminal, and tort law, statutory interpretation, international law, and procedure. It is the most comprehensive public measure of legal reasoning in language models and exercises exactly the capabilities the domain adaptation targets.
- **CUAD** [3]: commercial contract clause extraction across 41 clause types (see the Contract Analysis section above).
- **ECHR judgment prediction** [4]: legal judgment prediction over European Court of Human Rights cases.
- **ContractNLI** [5]: document-level natural-language inference over contracts (e.g. obligation and entailment detection).
- **Multi-LexSum** [6]: abstractive summarization of civil-rights litigation records.

5.2 Reporting Protocol

For each benchmark we recommend disclosing: the exact model version and quantization, the prompt template, the decoding parameters, the evaluation harness, and a side-by-side result for the unmodified Qwen3-8B base. We also recommend reporting calibration (how often a stated confidence matches empirical accuracy) and citation-verifiability rates, which matter more for safe legal use than raw task accuracy. Case-outcome *prediction* is deliberately excluded from our recommended deployment evaluations: predicting how a court will rule is ethically fraught, prone to spurious correlation, and not a use we endorse for an augmentation tool.

6 Legal Research Workflow

6.1 Memorandum of Law Generation

Zen Legal generates complete legal memoranda following standard structure:

1. **Question Presented:** Precise formulation of the legal question.
2. **Brief Answer:** Concise answer with legal conclusion and confidence.
3. **Statement of Facts:** Relevant factual background.
4. **Discussion:** IRAC (Issue, Rule, Application, Conclusion) analysis with full citations.
5. **Conclusion:** Summary of findings and recommended next steps.

The appropriate evaluation of generated memoranda is a blind, attorney-graded comparison against human-drafted work product, scored on legal accuracy, citation completeness, analytical clarity, and conclusion soundness. We do not report scores from such a study in this report; any quality claim of this kind should be backed by a pre-registered evaluation with disclosed rubric, sample size, and inter-rater agreement. As a structured-generation tool, Zen Legal’s intended advantage is in citation completeness and systematic IRAC coverage rather than in replacing attorney judgment.

7 Professional Responsibility Compliance

7.1 Unauthorized Practice of Law

Zen Legal explicitly frames all outputs as legal research assistance, not legal advice. The system:

- Appends standard disclaimers: “This output does not constitute legal advice and does not create an attorney-client relationship.”
- Monitors for jurisdiction-specific UPL rules that constrain AI-generated legal content.
- For consumer-facing deployments, routes complex matters to licensed attorney review before delivery.

7.2 Competence and Supervision

Per ABA Model Rule 5.3, attorneys using AI tools must supervise the AI’s work competently. Zen Legal supports this by:

- Providing explicit confidence scores on all legal conclusions.
- Flagging recently decided cases (<6 months) that may not be fully integrated in training data.
- Recommending verification steps for time-sensitive or high-stakes matters.

8 Conclusion

Zen Legal demonstrates that rigorous citation grounding, jurisdictional awareness, and professional responsibility compliance can be layered onto an openly licensed general-purpose model through targeted domain fine-tuning. Built as an Apache-2.0 fine-tune of Qwen3-8B rather than a from-scratch system, it inherits a strong, transparently licensed foundation and adds the retrieval, verification, and guardrail machinery that legal use demands. We have intentionally avoided headline benchmark claims in this report; the honest measure of Zen Legal is a reproducible, attorney-reviewed evaluation against public benchmarks such as LegalBench and CUAD, reported side-by-side with the unmodified Qwen3-8B base so that the contribution of the legal adaptation is transparent.

References

- [1] Qwen Team, Alibaba Group. Qwen3 Technical Report. 2025. Models released under the Apache-2.0 license. <https://github.com/QwenLM/Qwen3>
- [2] Guha, N. et al. LegalBench: A Collaboratively Built Benchmark for Measuring Legal Reasoning in Large Language Models. *NeurIPS*, 2023.
- [3] Hendrycks, D. et al. CUAD: An Expert-Annotated NLP Dataset for Legal Contract Review. *NeurIPS*, 2021.
- [4] Chalkidis, I. et al. Neural Legal Judgment Prediction in English. *ACL*, 2019.
- [5] Koreeda, Y. and Manning, C. ContractNLI: A Dataset for Document-level Natural Language Inference for Contracts. *EMNLP Findings*, 2021.
- [6] Shen, Z. et al. Multi-LexSum: Real-World Summaries of Civil Rights Lawsuits. *NeurIPS*, 2022.