

TEXAS ENVIRONMENTAL LAW SUPERCONFERENCE

AUGUST 7, 2026

# Environmental AI

How LLMs are rewiring environmental and administrative practice

**Tracy Hester**

Professor, University of Houston Law Center



# Roadmap

- Get oriented: some introductory key terms and concepts
- Close to home: impacts on legal education and training
- Looking ahead: how environmental law could change in response



# The Growing Field



## CHATBOTS

Conversation-first AI

EXAMPLES\*



ChatGPT.com



Claude.ai



Kimi.com

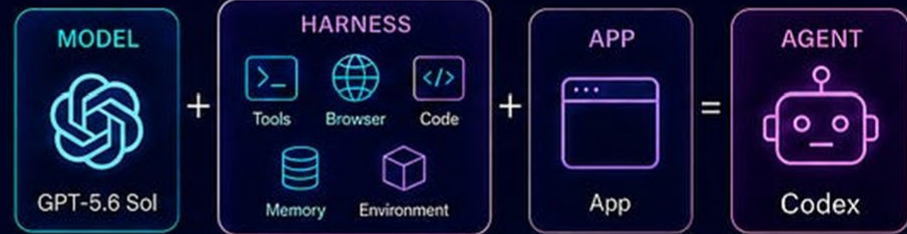
- You ask, it responds
- Best for Q&A, brainstorming, drafting, and short tasks
- Usually turn-by-turn and user-directed
- Limited autonomy



## AGENTIC SYSTEMS

Goal-directed AI systems

MODEL + HARNESS + APP = AGENT



Codex = GPT-5.6 Sol + a harness + an app

EXAMPLES

- You give a goal, not step-by-step instructions
- Plans and acts over longer tasks
- Uses tools, memory, and external systems
- You supervise and evaluate results

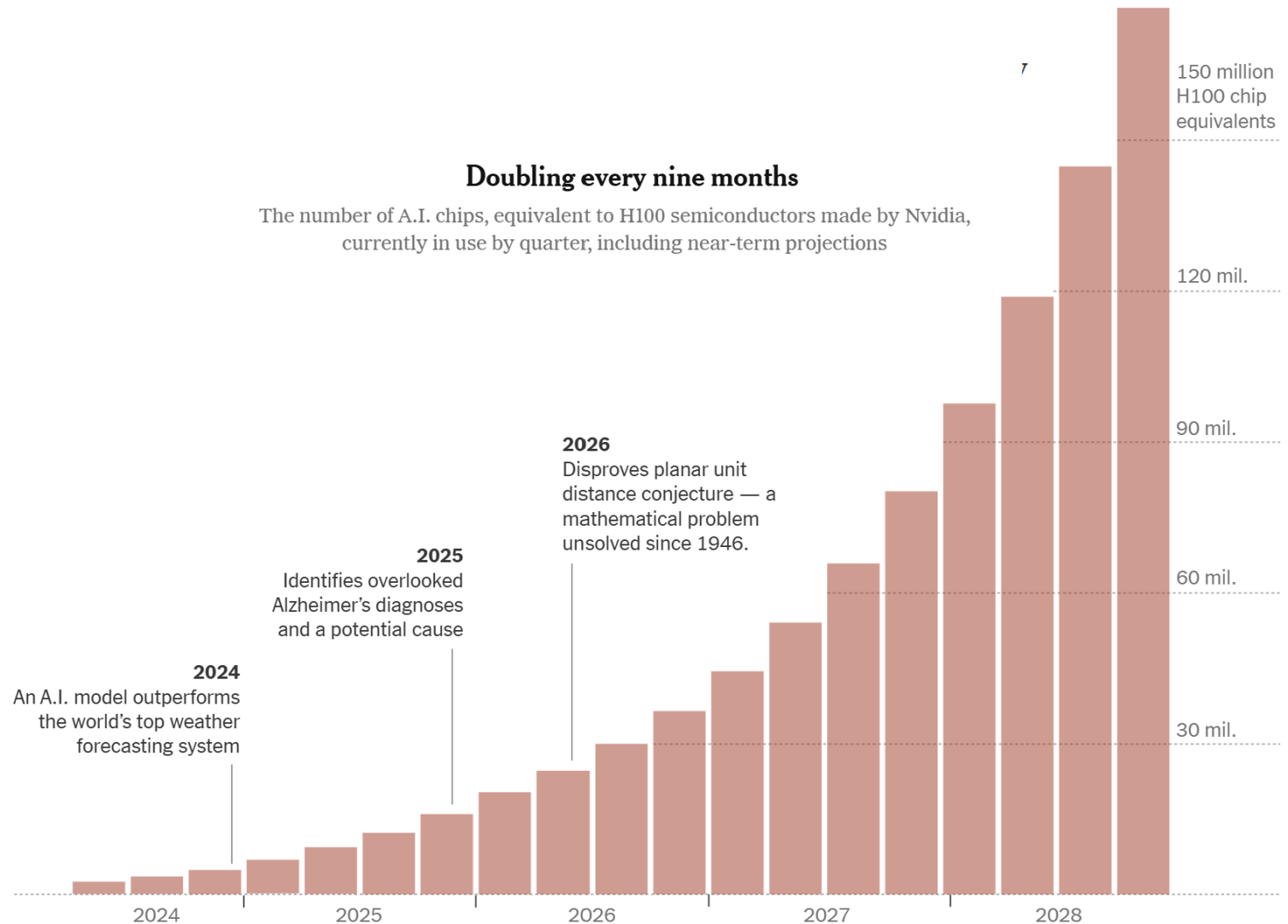
- ✦ Gemini Notebook
- ✦ Claude Code
- ✦ GPT Work
- ✦ Manus
- ✦ Grok Build

\* You can also access agentic systems through these sites

oneusefulthing.

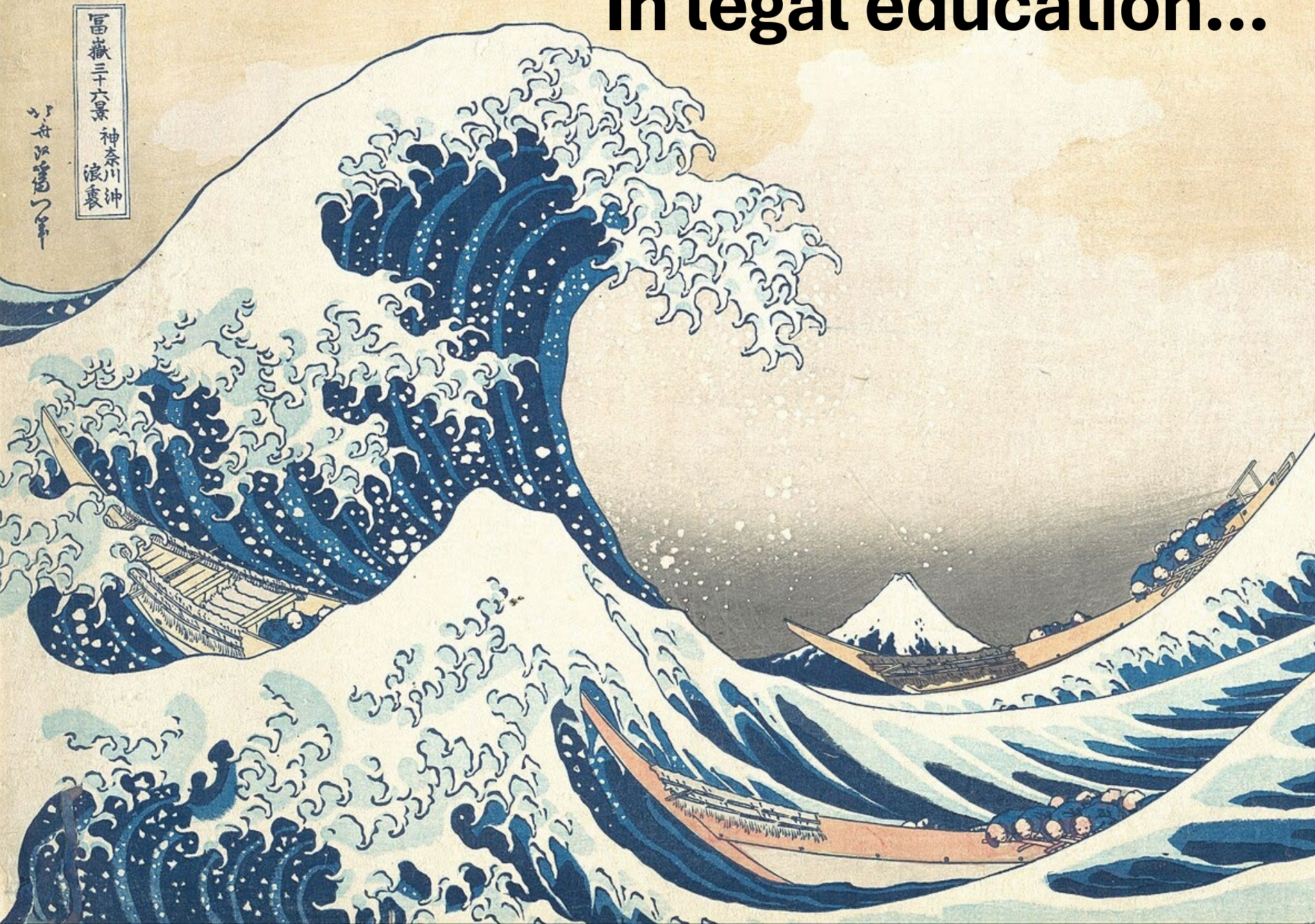
|                                                                                                                                                    |  MODEL NAME                                                         |  THINKING / EFFORT LEVEL |  BEST USE / NOTES              |
|----------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------|-----------------------------------------------------------------------------------------------------------|-------------------------------------------------------------------------------------------------------------------|
| <br><b>THE VERY SMARTEST MODELS</b><br><i>(\$20/month+)</i>       | GPT-5.6 Sol                                                                                                                                          | Pro                                                                                                       | Best for hard math or science                                                                                     |
|                                                                                                                                                    | Claude Fable 5                                                                                                                                       | High, Extra or Max                                                                                        | Can work independently for the longest                                                                            |
|                                                                                                                                                    | GPT-5.6 Sol                                                                                                                                          | Extra High or High                                                                                        | All around good model, a little below Fable                                                                       |
| <br><b>SMART MODELS</b><br><i>(usually requires subscription)</i> | Kimi K3                                                                                                                                              | Max or High                                                                                               | The strongest open weights models                                                                                 |
|                                                                                                                                                    | Claude Opus 4.8                                                                                                                                      | Max or High                                                                                               | Less weird (and less powerful) than Fable                                                                         |
|                                                                                                                                                    | GPT-5.6 Sol                                                                                                                                          | Medium                                                                                                    | Faster, slightly less good answers                                                                                |
| <br><b>BEST FREE MODELS</b>                                      | <ul style="list-style-type: none"> <li>• Gemini 3.6 Flash</li> <li>• GPT-5.5 Instant</li> <li>• Claude Sonnet 5</li> <li>• Muse Spark 1.1</li> </ul> | N/A                                                                                                       | These are models from the US  |
|                                                                                                                                                    | <ul style="list-style-type: none"> <li>• DeepSeek V4 Preview</li> <li>• Qwen 3.6 Plus</li> </ul>                                                     | N/A                                                                                                       | These are models from China  |

# Navigating the Tricky Pitfalls of the Wait Calculation...



Note: The data here represents H100 equivalents as measured by chips in use. Source: Epoch AI.

**In legal education...**



## INITIAL IMPACTS

# Law Schools - Preparing New Lawyers for AI

*Two elite schools set two different paths, six weeks apart — and neither is where new hires will actually get their training.*

**3%**

of law schools require any AI coursework

**2.44 / 5**

faculty confidence in their students' AI skills

**19.6% → 68.7%**

student confidence, before vs. after their first work experience

### Protect the formative year

**Chicago, July 2026.** Laptops banned in every 1L classroom, closed-book in-class exams, AI built deliberately into legal research and writing, supervised AI in the clinics. The goal is a graduate who can think **with, without, and about** AI. The premise: a 1L's ability to judge whether model output is any good is weakest exactly when the temptation to lean on it is strongest.

### Level access up, and define competence

**Texas, June 2026.** Enterprise licenses with frontier labs for every student, plus Harvey and Legora, so tool access stops being a class advantage. Take-home exams nearly eliminated in favor of locked-down in-class exams, live presentations, and some oral exams. And in writing seminars, faculty are beginning to require students to **document their AI use step by step — and to be graded on the quality of it.**

**Effects on jobs and skills expectations-** UT's learning objectives ask for graduates who can produce verified work product, justify choosing one tool over another on capability, risk and cost, and know where AI use touches privilege and work product. But the value of a polished writing sample, obviously, is much degraded..

# KEY IMPACTS OF AI ON LEGAL WORK

**Producing sophisticated legal work product historically has been expensive, bespoke, and time-consuming. That cost wasn't a restriction or cap on legal work, but it worked like one.**

It kept marginal claims out of court and made clients ration their lawyers' time. And it put many categories of analysis — surveying every permit term in a program, reconstructing what a technical regulatory term meant in 1976, stress-testing a rule against every possible exploit — out of practical reach.

# KEY IMPACTS OF AI ON LEGAL WORK

1

## The same work, faster

20–28% less time and +0.25–0.53 quality points across six realistic tasks — with the largest gains going to the lowest-baseline performers.

*Schwarcz et al., first RCT (2026)*

2

## More people use it

Pro se filings account for 59% of all growth in non-prisoner federal civil litigation since 2023.

*Shah & Levy (2026)*

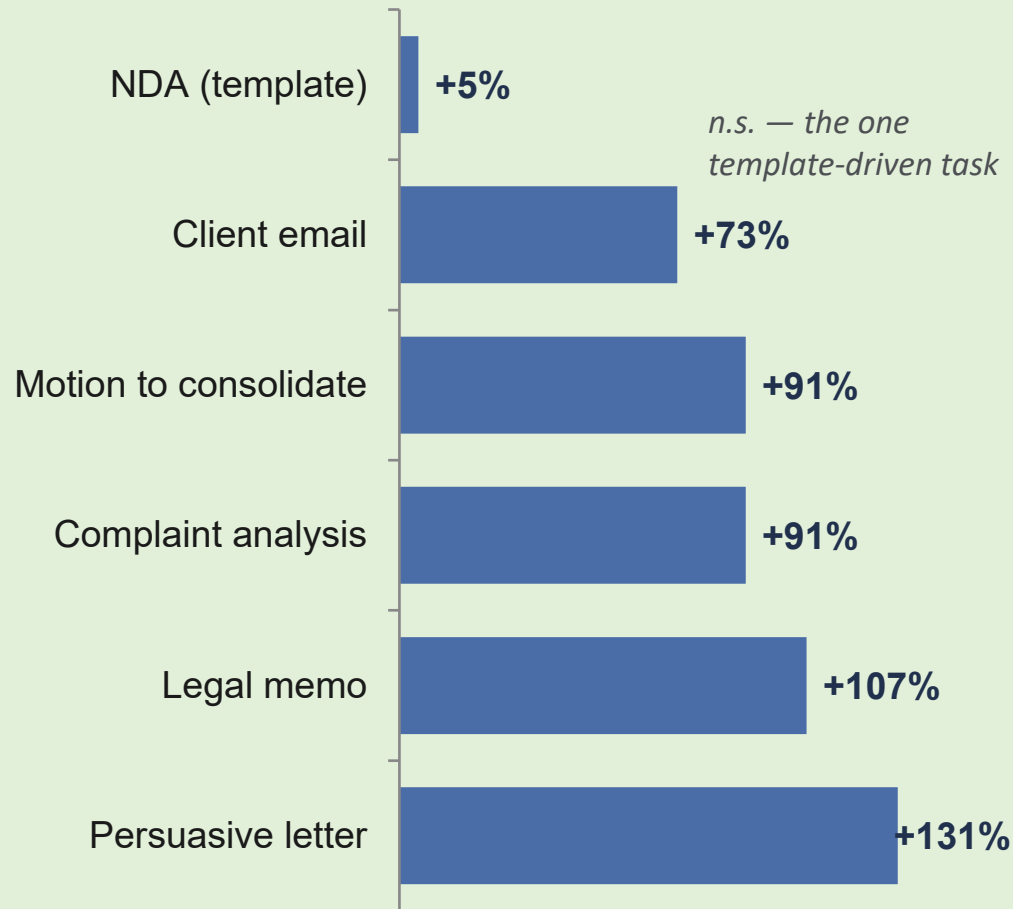
3

## Work that was impossible

Whole-corpus analysis, quantified ambiguity, adversarial audits of rules that have not issued yet.

## LOWERED BARRIERS

# More work, done more quickly



**20–28%**

less time on 5 of 6 tasks

**+0.25–0.53**

quality points (7-pt scale, blind-graded)

**3 · 4 · 11**

hallucinated cites: RAG · no-AI · reasoning model. Grounding beat working alone.

# Twenty years flat *pro se*, then 2023



## PRO SE SHARE OF NON-PRISONER FEDERAL CIVIL FILINGS

*Two decades at ~11% — unmoved by the Great Recession or the ACA wave.*

**59%**

of all growth in non-prisoner federal civil filings is new *pro se* cases

**~18%**

of 2026 complaints flagged as AI-drafted — from roughly zero before 2023

### Why this is a challenge.

This surge concentrates in simple, document-driven case types — and disproportionately in suits against federal agencies or institutional defendants.

# Large Language Models Hack Rewards, and Society

Wei Liu<sup>★\*✉</sup> Xinyi Mou<sup>♠\*</sup> Hangi Yan<sup>★</sup> Zhongyu Wei<sup>♠♥</sup> Yulan He<sup>★♣✉</sup>

ARTIFICIAL INTELLIGENCE

## AI finds legal loopholes, even when not asked

Models repeatedly discovered ways to exploit regulations,  
and current safeguards largely failed to stop them

**Correspondence:** {wei.4.liu, yulan.he}@kcl.ac.uk

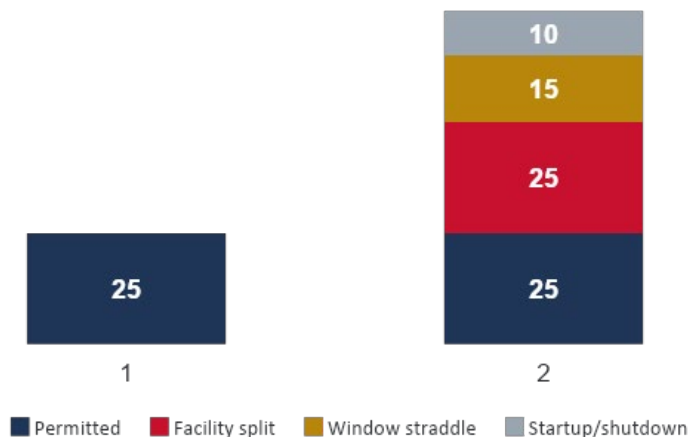
**Code:** <https://github.com/thinkwee/SocioHack>



# Simulation: stress-testing rules before they issue

**THE FINDING** Across 72 simulated regulatory environments, an LLM rediscovered more than 60% of real, later-patched loopholes — including one reform not yet enacted — and self-flagged only 37% of its own exploitation.

**Applied to a draft rule:** hypothetical TCEQ permit-by-rule — 25 tpy VOC per “facility,” rolling 12-month average, with a startup/shutdown exemption. Formalize the rule → let an adversarial model maximize throughput within its literal text → harvest the exploits → attach a patch to each.



The rule caps 25 tons. Its literal text permits 75. **Three times the cap.**

## Suggested fixes

- Define “facility” by contiguous site plus common control.
- Add a fixed calendar-year backstop to the rolling window.
- Cap exempt-event hours per quarter.

# Empirical statutory interpretation

- “Ordinary person” textualism and LLMs
- Time machine originalism and curated LLMs
- Testing of regulatory interpretations – RCRA reuse, Clean Air BACT

## Generative Interpretation

*Yonathan Arbel & David A. Hoffman\**

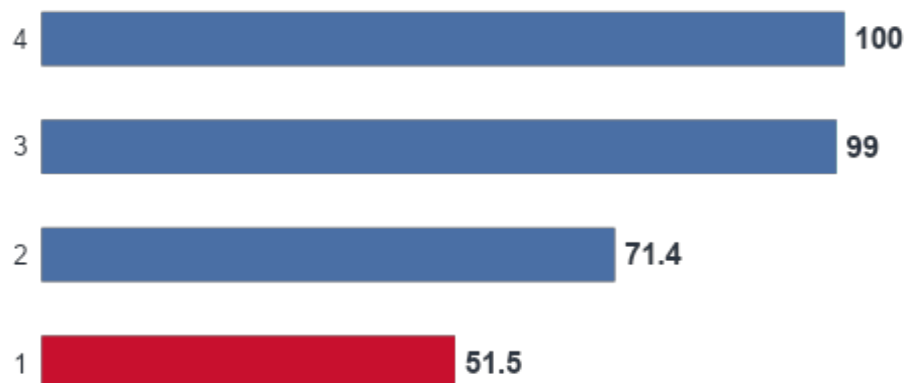
[*DRAFT*] July 30, 2023]

*We introduce generative interpretation, a new approach to estimating contractual meaning using large language models. As AI triumphalism is the order of the day, we proceed by way of grounded case studies, each illustrating the capabilities of these novel tools in distinct ways. Taking well-known contracts opinions, and sourcing the actual agreements that they adjudicated, we show that AI models can help factfinders ascertain ordinary meaning in context, quantify ambiguity, and fill gaps in parties' agreements. We also illustrate how models can calculate the probative value of individual pieces of extrinsic evidence.*

# “Solid waste” under RCRA, run 400 times

**QUERY** Is K051 API separator sludge sent for reclamation “discarded material,” and thus “solid waste” under RCRA § 1004(27), in ordinary usage?

SHARE ANSWERING “NOT DISCARDED,” BY HOW THE QUESTION WAS ASKED



**80.5%**

of all 400 readings said “not discarded”

**49 points**

of it moved on nothing but the wording of the question

**The results reflect the method, rather than the word itself.** Ask a LLM neutrally and ordinary usage splits nearly evenly. Supply a persona and it goes unanimous. Gives reason to run the simulation hundreds of times across framings rather than once; be careful to avoid the inherent bias in a single one-off prompt.

ALREADY DEPLOYED

# Agencies now use AI in environmental reviews

*Note: These uses are not hypothetical – already named, funded, operating programs, and many more will likely to follow.*

## PermitAI

DOE / PACIFIC NORTHWEST NATIONAL LAB

Purpose-built to accelerate environmental review and permitting by distilling decades of prior NEPA documents.

## NEPA drafting

DOI · FEDERAL PERMITTING COUNCIL

AI used to draft and assemble environmental review documents and to analyze underlying technical data.

## State programs

MINNESOTA · CALIFORNIA

State agencies applying the same tooling to their own environmental review and permitting workloads.

## Comment triage

REGULATIONS.GOV / ICF

Generative models used to process and summarize public comment volumes that no staff could read.

### For practitioners, note four risks:

- Hallucinated content inside an administrative record;
- A model whose reasoning cannot be reconstructed for judicial review;
- Uncertain judicial treatment of machine-assisted findings; and
- An administrative record that may no longer contain everything the agency actually relied on.

# Optimizing comments for AI readers

## AVOID

- X Burying key points in body of text
- X Relying on nuance that AI models will overgeneralize away
- X Arguments that only work by cross-reference
- X Labeling “Additional Comments” as a header

## EMPHASIZE

- ✓ Use substantive headers that state the issue
- ✓ Emphasize declarative topic sentences, one point per section
- ✓ Pick explicit citations — read as relevance signals
- ✓ Employ strategic repetition: summary → body → conclusion

Summarizers weight keyword frequency, sentence position, and semantic context — and overgeneralize nuance. Substance still wins; structure decides whether it is seen.

# The APA and Agency AI Use

*Existing principles of judicial review for environmental-rulemaking authority already apply to AI actions.*

## Disclosure

### *Portland Cement Ass'n v. Ruckelshaus*

An agency must reveal the critical factual material, data, models, and assumptions it relied on. If the “model” is a proprietary LLM with undisclosed training data and unlogged prompts, what exactly has been disclosed?

## Reasoned decisionmaking

### *State Farm · Ohio v. EPA (2024)*

The agency must consider and respond to significant comments and explain its reasoning. An agency that rubber-stamps what a model chose has not deliberated — and first drafts are sticky.

## Likely judicial expectations from agencies

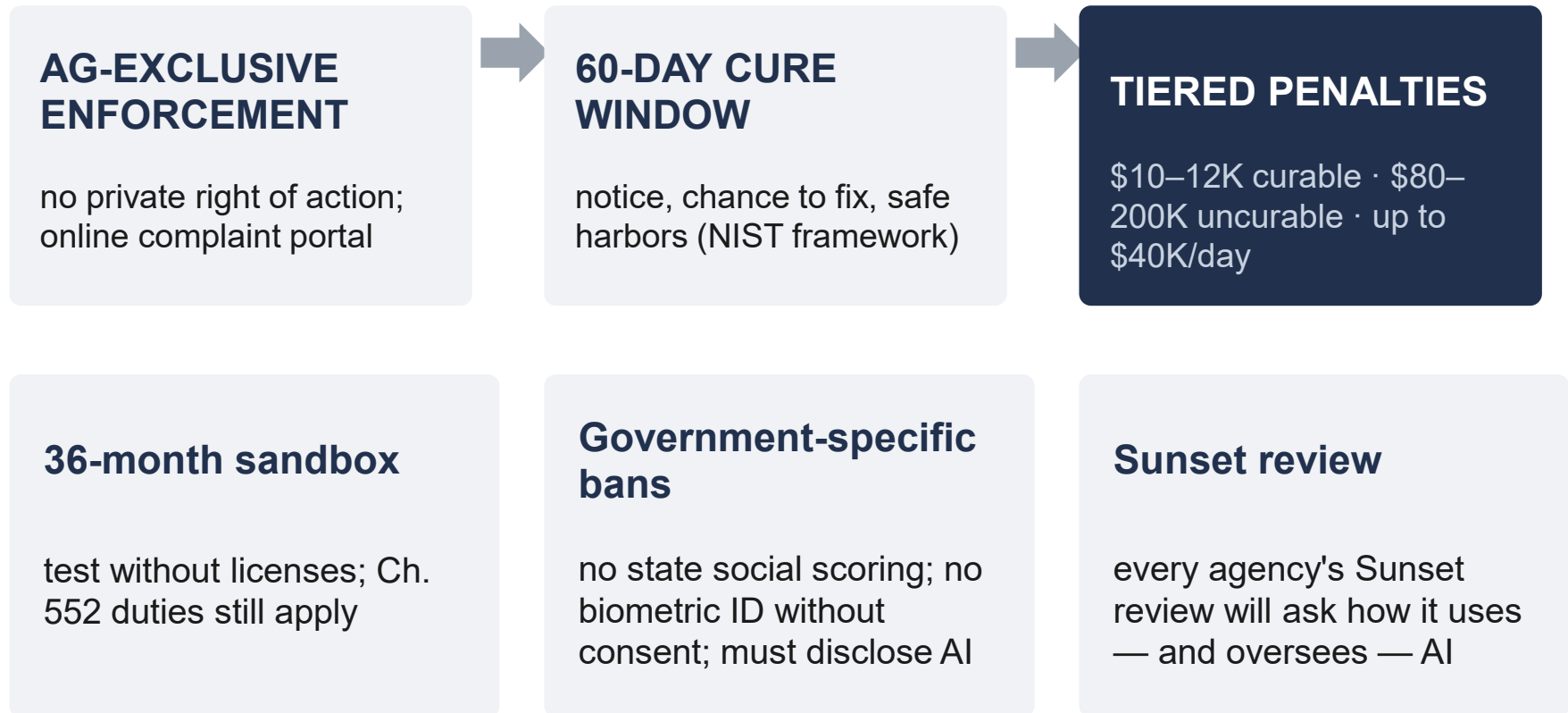
*Seven questions, none of them exotic, all of them answerable:*

- Which model, at which version?
- Procured how, and under what terms?
- Fine-tuned — and on what data?
- What prompts were used?
- What raw outputs were produced?
- What quality control and validation?
- Who outside the agency was involved?

# TRAIGA and Environmental Law

(Texas Responsible Artificial Intelligence Governance Act)

*Requires transparency for algorithmic permit processing, bans intentional discrimination, and data integrity requirements for machine learning uses*



**H.B. 149 (TRAIGA), 89th Leg. — effective January 1, 2026**

# Beware:

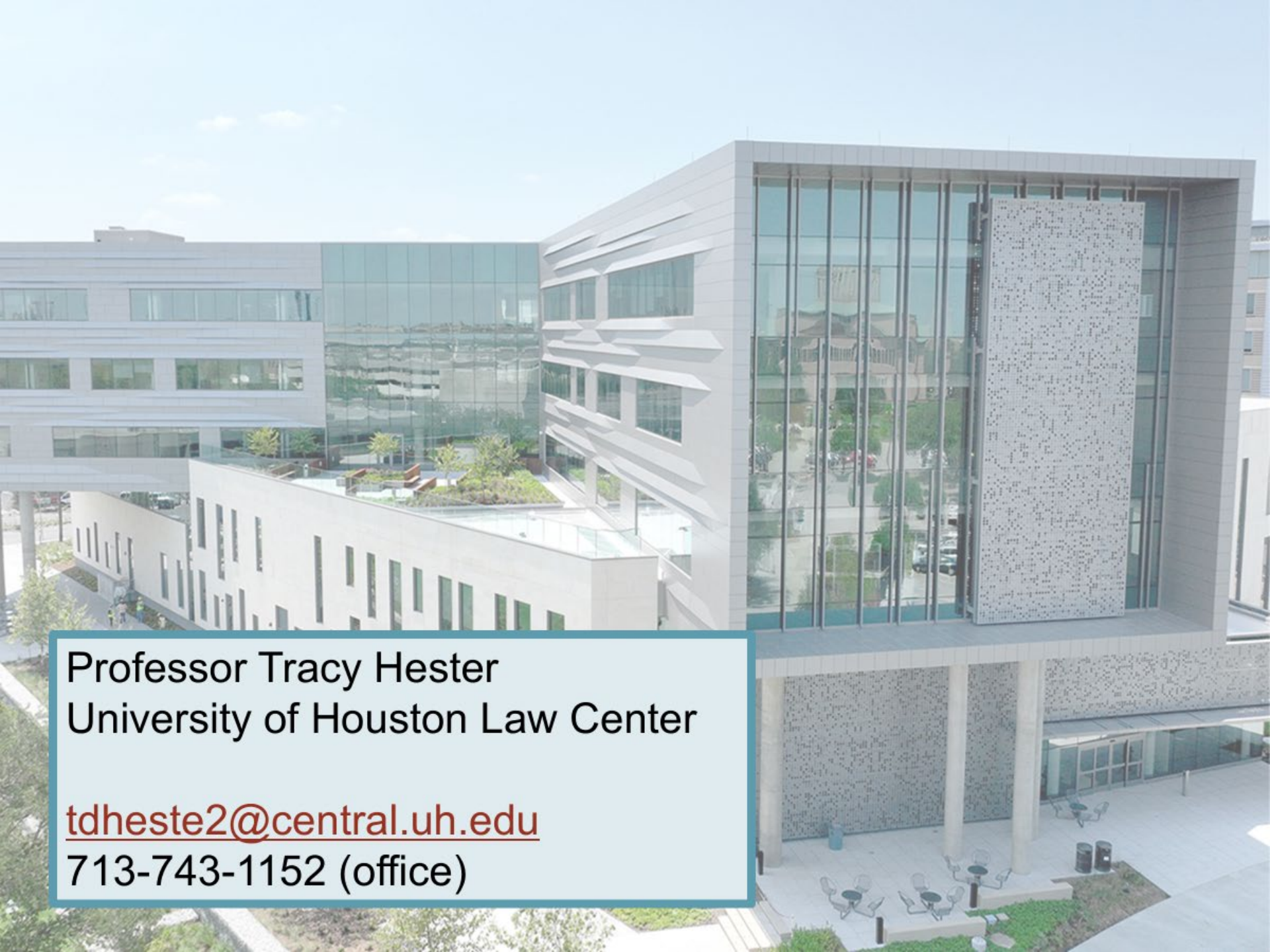
- Biases of synthetic persons used in simulations
- Ethical obligations in court and agency representations
- Security vulnerabilities
- Expert opinions: prompts, communications, bounded datasets





# Next Horizons

- Self-directed AI legal research with autogenerated hypotheses
- Autonomous corporate clients
- AI native firms
- AI arbitrations



Professor Tracy Hester  
University of Houston Law Center

[tdheste2@central.uh.edu](mailto:tdheste2@central.uh.edu)  
713-743-1152 (office)