

Quarterly Research Update

Q2 2026 | April–June 2026

This update highlights research and practical resources released—or newly surfaced and not previously covered—during Q2 2026. The strongest materials examine how technology may expand access and efficiency while procedural fairness still depends on meaningful participation, understandable processes, neutral review, respectful treatment, and visible human responsibility.

Top takeaways

- **AI is changing both sides of the courthouse threshold.** Litigants are increasingly able to prepare polished filings without lawyers, while courts are beginning to test AI tools to support review of high-volume matters. The procedural-justice question is not simply whether AI is used, but whether it improves meaningful voice, understanding, neutrality, and trustworthy decision-making.
- **Formal participation is not necessarily meaningful participation.** A filing may look legally sophisticated without improving the filer's understanding or prospects. Likewise, an absent defendant receives meaningful judicial review only if the court gives a default request careful and genuinely independent scrutiny.
- **Human responsibility and differing community perspectives must remain visible.** AI tools can make review and analysis more efficient, but they should not obscure who makes the decision, what assumptions shape the result, how respect is understood by different groups, or how a person can challenge an error.

Featured research: Can AI improve review of default judgments?

AI Assistance for Human Review of Default Judgments (preprint, submitted June 4, 2026) examines a setting in which procedural safeguards are especially important. In debt-collection default cases, the defendant often does not appear, leaving the court responsible for determining whether the plaintiff has satisfied the legal requirements for judgment. Because the proceeding is effectively one-sided, careful and neutral court review becomes a substitute for the adversarial testing that would otherwise expose errors.

The researchers audited 188 Los Angeles Superior Court debt-buyer cases in which default judgment was ultimately granted. They found that 4% contained major defects that should have prevented judgment, 10% contained inconsistencies requiring a lower judgment, and 32% contained errors requiring amendment before judgment.

The paper estimates that court research attorneys average about four minutes per case to verify more than a dozen statutory requirements. The researchers then tested a “Default Assistant” that gives reviewers requirement-level recommendations supported by quotations and tables drawn from the actual filings. Sixty-six law students reviewed cases in a controlled study. They were allowed up to ten minutes per case—still a modest period for reviewing a request that can result in a binding judgment, garnishment, and

other serious consequences. Students using the assistant were 6.0% more accurate on the average requirement and 25.9% faster than students reviewing the same materials without AI assistance.

The contrast is striking, but this was not a direct experiment comparing AI-assisted students with court attorneys, and the students received an explicit checklist that court staff apparently do not ordinarily use. Even so, the results suggest that a checklist, adequate review time, and verifiable AI assistance could help courts identify defects now being missed.

The paper also leaves an important accountability question unanswered. It says research attorneys conduct the initial review and that final decisions remain with staff attorneys and judges, but it does not explain how much independent review the responsible judges performed or how recommendations moved from staff to judges. All 188 audited requests were ultimately granted. Because entry of judgment is a judicial act, reforms should clarify not only what the technology and court staff review, but also what the judge personally examines and decides.

This remains a proof of concept involving law students, one jurisdiction, and one case type; field testing with court personnel is still needed. But when one party is absent and has little effective voice, neutral, transparent, and careful court review is indispensable.

New research: AI-assisted self-representation

The New Pro Se: Generative AI and the Surge in Federal Civil Self-Representation (preprint, submitted May 28, 2026) provides large-scale evidence about a phenomenon judges and court staff are increasingly observing. The study analyzes approximately 2.8 million federal civil filings from fiscal years 2008 through 2025. It reports that the share of federal civil plaintiffs proceeding without counsel rose from 11.33% before widespread public access to generative AI to 16.94% afterward.

In a narrower complaint-text analysis, the author developed a stylometric measure of “AI-consistent drafting.” Against a pre-AI baseline, the estimated net AI-flagged share was 13.9% of post-AI non-form complaints. Those complaints were more citation-dense and disproportionately associated with people first observed as filers in the study sample. But they did not have better win rates; they were more likely to be dismissed and to end at earlier procedural stages.

The paper does not observe whether a particular litigant actually used AI, and it cannot prove that generative AI caused the overall increase. Its findings nevertheless provide systematic evidence concerning a change courts are already seeing: AI can help produce a document that looks like a legal pleading, but polished drafting is not the same as understanding the process, stating a sufficient claim, or having an effective opportunity to be heard.

A pleading may overstate the filer’s actual comprehension, making it harder to recognize when explanation or assistance is needed. Courts should preserve access without treating surface sophistication as proof of understanding. Plain-language instructions, self-help support, warnings about AI limitations, and opportunities to identify and correct deficiencies are more consistent with procedural justice than ignoring the phenomenon or responding mainly through punishment.

Practical guidance: Build safeguards around court use of AI

NCSC's *AI Foundations in the Courts* (April 23, 2026) supplies a practical counterpart to the two research papers. It advises courts to establish clear AI-use guidelines, maintain human oversight, test for bias, verify outputs, protect confidential information, and tell court users when AI may affect them. It also warns that machine translation cannot replace qualified interpreters for court events or complex legal communication.

The guidance reinforces a central lesson: keeping a human nominally “in the loop” is not enough. Courts should identify the responsible official, supply verifiable source material, monitor performance, and preserve meaningful ways to question or correct a result.

Practical implementation: Put community members inside the evaluation process

NCSC's May 19 case study, *Using Community Insights to Strengthen a Community Court in Georgia*, describes a participatory-action-research process used to evaluate Albany Works! Community Court. Instead of relying only on court-led research and administrative data, the project involved community partners, residents, students, and program participants in developing the research and interpreting the court's performance.

Participants rated the program an average of 4.3 out of 5 and highlighted fair treatment, clear expectations, accountability, and support. Yet the evaluation identified a major trust gap: only 11% of stakeholders viewed court-community relationships positively.

That contrast gives the study much of its value. Positive experiences inside one program do not automatically produce broader institutional trust. The participatory method is itself a procedural-justice practice because community members help define success and identify needed changes.

Practical implementation: Preserve voice and autonomy in guardianship

NCSC's June 10 case study, *Supporting Vulnerable Adults Through Community Guardianship Mapping in Massachusetts*, describes workshops that mapped community supports and less-restrictive alternatives to guardianship. The procedural-justice connection is direct: because guardianship can remove substantial decision-making authority, a fair process should promote informed participation, preserve autonomy where possible, and tailor authority to actual need.

Also of note: Hidden assumptions in automated bail tools

Confronting Label Indeterminacy in Automated Bail Decisions, submitted April 13, 2026 and presented at the June 2026 International Conference on Artificial Intelligence and Law, examines a basic problem in predictive bail systems. When a person is detained, no one can observe whether that person would have appeared in court if released. Historical data therefore contain an unknowable counterfactual.

Different ways of filling that gap can change predictions as much as, or more than, the model selected. Apparently technical choices may therefore determine who is labeled risky. Neutrality and legitimacy require those assumptions to be visible, challengeable, and recognized as policy choices rather than objective facts.

Earlier publication worth noting: Who decides what counts as respect?

The Subjectivity of Respect in Police Traffic Stops: Modeling Community Perspectives in Body-Worn Camera Footage, first posted in February 2026 and newly reviewed for this update, examines one of procedural justice's central dimensions: respectful treatment. The researchers developed a large dataset from roughly 1,000 LAPD traffic stops and asked annotators from three groups—people with law-enforcement experience, people with histories of arrest or incarceration, and other Los Angeles residents—to rate and explain the respect shown by officers and drivers.

The project rejected a single objective “ground truth” about respect. Its rubric drew on procedural-justice theory, LAPD materials, focus groups, interviews, a survey of more than 2,000 residents, and other fieldwork. An annotator's relationship to the justice system was the most important background factor associated with judgments of respect; age, race, and gender did not predict ratings in the same way.

Its broader contribution is methodological: agencies that evaluate interactions only through an institutional lens may miss how the same conduct is experienced by people with different histories. AI analysis of body-camera material may need to preserve multiple perspectives rather than collapse them into one official definition of respect. Because this study uses transcripts rather than tone, facial expression, or body language, it is not a complete measure of an encounter.

Suggested actions

- **Audit default-judgment review using the study's framework.** Measure actual time per case, document the staff-to-judge workflow, identify who checks each statutory requirement, and clarify the extent of the judge's independent review. Test checklists and citation-grounded assistive tools with court personnel before deployment.
- **Prepare court staff for AI-assisted self-representation.** Using the problems identified in *The New Pro Se* and the safeguards in NCSC's *AI Foundations in the Courts*, develop guidance for clerks, self-help personnel, and judges that distinguishes polished drafting from actual understanding and emphasizes explanation, correction, and meaningful participation.
- **Build user and community voice into evaluation.** Adapt the Albany participatory-evaluation model by involving court users and community stakeholders in defining success, interpreting results, and identifying needed changes—not merely responding to a finished court-designed survey.
- **Make the assumptions behind automated decisions visible and contestable.** For any AI or risk-assessment tool, document what unobservable outcomes have been assigned labels, who selected those assumptions, how alternative assumptions change the result, and which human official remains responsible for the final decision.

The Procedural Justice Hub is a service of the Procedural Justice Project, founded by Emily LaGratta, Judge Kevin Burke, and Steve Leben. | proceduraljustice.org