

Access to Justice in the Age of AI: Evidence from U.S. Federal Courts*

Anand V. Shah[†]

Joshua Y. Levy[‡]

Massachusetts Institute of Technology

University of Southern California

April 2026

Abstract

This paper studies how generative AI has reshaped entry into the federal civil court system. Drawing on administrative records covering more than 4.5 million non-prisoner federal civil court cases from FY2005-FY2026 and 46 million PACER docket entries matched to those cases, we document three sets of findings. First, the number of *pro se* cases—or self-represented cases—is increasing dramatically, rising from a long-term steady-state average of 11% to 16.8% in FY2025. This increase is concentrated in case types characterized by formulaic document production and absent from more complex, attorney-intensive categories. Second, we argue these cases are placing a larger burden on federal district courts. *Pro se* cases are not terminating faster, and this combined with the increased case numbers suggests more cases for judges to process. Moreover, intra-case activity is up, with the total volume of docket entries per court generated by *pro se* cases in their first 180 days up 158% from pre-AI means to 2025. Third, we directly validate that AI use is increasing in federal courts. Using a random sample of 1,600 complaints drawn from an 8-year period (2019-2026), we find that a large and growing share of complaints are flagging positive for AI-generated text, from essentially zero in the pre-AI period to more than 18% in 2026.

Keywords: access to justice, generative AI, pro se litigation, legal technology

*We'd like to thank Elliott Ash, Bao Kham Chau, Basil Halperin, Brian Jabarian, Sara Jain, David Schönholzer, Sankalp Sharma, Dan Sokol, Jordi Weinstock, and Parker Whitfill for helpful discussion. We are also grateful to Pangram Labs for access, and to MIT's Initiative on the Digital Economy for support. All errors are our own. JEL Codes: K40, J24, O33, L86.

[†]avshah@mit.edu

[‡]jylevy@usc.edu

It's easy with these tools to churn out hundreds of thousands of works in the time that a human author would produce maybe one or two... what we have is a room of screaming toddlers, and we can't hear the people we're trying to listen to.

— Feb. 21, 2023, Neil Clarke, editor of *Clarkesworld* on a flood of LLM-generated submissions.

1. Introduction

For those in the business of science fiction, the five-time Hugo Award-winning magazine *Clarkesworld* is well known as one of the most prestigious literary venues. Yet, in February 2023, the magazine faced a crisis. Inundated with quintuple the typical submission count, the magazine chose to temporarily shut down, unable to scale its editorial capacity to the increased submissions attributed to unprecedentedly powerful LLMs. Three years and many models later, we find ourselves still asking how best to adapt to the decreased cost of generating adversarially passable information through our evolved filtering mechanisms.

Yet, while magazines can hire more editors, there is no easy margin along which to “buy” extra judge capacity. Already case backlog is becoming a persistent feature of the federal judicial system, there is no coming influx of judges to supply additional capacity,¹ and federal courts in the United States cannot wholesale decline to hear cases (Menell and Vacca, 2021; Administrative Office of the U.S. Courts, 2024). If generative AI dramatically lowers the cost of self-represented litigation, the resulting surge in filings could overwhelm a system that depends on human judgment at every stage of adjudication. This paper examines whether that glut has arrived and, if so, how the courts are digesting it.

A reasonable way to measure the cost of entry into courts is along people choosing to file *pro se*. The right to represent oneself in legal proceedings, or *pro se* litigation, is among the oldest statutory rights in American law, older even than the Bill of Rights.²³ For most of our

¹Only an act of Congress can expand the number of Article III judges, and the number of Article III judge appointments has grown only modestly for decades.

²The right to appear *pro se* in federal civil cases was first enacted as Section 35 of the Judiciary Act of 1789, signed into law on September 24, 1789—one day before the Bill of Rights was proposed to the states.

³Will Hunting: “I am afforded the right to speak in my own defence, sir, by the Constitution of the United States! This is the same document which guarantees my liberty. And liberty, in case you’ve forgotten, is a soul’s right to breathe. And when it cannot take a long breath, laws are girded too tight. Without liberty, man is a syncope.” Judge: “Man is a what?” Will Hunting: “Ibid, Your Honour.” – *Good Will Hunting*

courts’ history, this right has been exercised at a “remarkably unremarkable” steady state. Gough and Taylor Poppe (2020) examine federal civil filings from 1999 to 2018 and document a civil litigation *pro se* representation rate of approximately 11%—stable across two decades and notably unperturbed by the Great Recession, the Affordable Care Act litigation wave, or the surge in immigration cases under the first Trump administration.

This stability seems to reflect a structural barrier: for most people, self-representation is prohibitively hard. Filing a federal civil complaint requires identifying the correct jurisdictional basis, pleading sufficient facts to survive a motion to dismiss, and navigating procedural requirements that vary by context and case type.

The widespread, public diffusion of capable LLMs changes that calculus. Without a law degree and at *de minimis* cost, any person with an internet connection can not only obtain interactive, case-specific legal guidance—drafting complaints, identifying statutes, navigating procedure—but also generate passable legal documents, particularly so after the release of GPT-4 in March 2023.⁴ Whether these technological advances have been sufficient to materially alter filing behavior is an important empirical question for the future of our courts. We study it using administrative records covering roughly 4.5 million non-prisoner federal civil cases filed from fiscal years 2005-2026. We draw on two complementary datasets: filing-based counts from the Administrative Office of the U.S. Courts (AOUSC) and the Federal Judicial Center Integrated Database (IDB), which tracks filings, dispositions, and durations at the case level. We supplement the IDB with 46 million PACER docket metadata entries that enable us to look “inside” cases to measure motion activity.

We document three sets of findings.

First, the national non-prisoner *pro se* filing share rose sharply from its approximately 11% historical steady state to 16.8% in fiscal year 2025, a gain that has no precedent in 25 years of administrative records. Moreover, that rise is not uniform across case types. It is concentrated in categories where the dominant task is producing a well-structured factual narrative—civil rights complaints, consumer credit disputes, foreclosure proceedings—and is essentially absent

⁴Katz et al. (2024) reported that GPT-4 scored approximately 297 on a full simulated Uniform Bar Examination, above the passing threshold in every jurisdiction and at roughly the 75th percentile of February 2022 test-takers. Guha et al. (2023) found that GPT-4 was the best-performing commercial model on LegalBench at release, with category averages of 83% on issue spotting, 90% on rule conclusion, 75% on interpretation, and 79% on rhetorical understanding. Models from many different providers have continued to push capabilities forward.

in categories that require sustained expertise, such as patent infringement and securities fraud.

Second, cases are not resolving any faster and analysis of case disposition suggests little change in the quality of cases brought by *pro se* litigants. However, more is happening within each case. Since the advent of capable LLMs, the total volume of *pro se* docket entries per court in the first 180 days of a case has grown by 64% on average across the post-AI period (Q4 2022–Q2 2025), and by 2025Q2—the most recent fully observable quarter at the 180-day horizon—this volume has more than doubled (+158%) the pre-AI mean. Each docket entry is a claim on the court’s time—a motion filed, a response submitted, an order entered, a hearing set. These results, bundled, suggest increased court burden in the post-AI period.

Third, we directly test that AI use is increasing in the post-period by applying an AI-content detector to a random sample of 1,600 federal civil complaints drawn from CourtListener’s RECAP archive. The AI-content detector records only one false positive (out of 800 queries) in the pre-period. We find that the share of complaints containing AI-generated text has been rising monotonically from 1.0% in 2023 to 18.0% in early 2026.

An important caveat is that this paper is necessarily descriptive, given that every internet user was given access to each LLM model at once, and that LLM use has diffused quite widely (see Figures 5 and 10). Moreover, differences in uptake are hard to justify without fine-grained usage data, given LLMs have seen the fastest adoption curve in the history of technology.⁵ The only counterfactual available is the pre-AI time series in each unit, and so the identification argument throughout this paper is at best temporal stability: the twenty years of *pro se* filing behavior that precede GPT-4 look the same as the ten years that precede it, so the post-2022 break is unlikely to be the continuation of a secular trend. We do not claim to identify a causal effect of GPT-4 on *pro se* filing, only that the observed time series is difficult to rationalize without generative AI playing a role. The visual tests in Figures 1, 4, and 5 are meant to make the pre-AI stability easy to see and to let the reader judge for themselves how unusual the post-LLM break is relative to earlier variation.

The paper proceeds as follows. Section 1.1 reviews the related literature and Section 2 describes the data. Section 3 documents the headline *pro se* results, characterizing who is

⁵<https://www.theguardian.com/technology/2023/feb/02/chatgpt-100-million-users-open-ai-fastest-growing-app>.

filing and in what case types. Section 4 examines the outcomes of these filings, including durations, dispositions, and intra-case docket activity. Section 5 reports the results of an AI-content detector applied directly to the complaint documents. Section 6 lays out a detailed agenda for future work. Section 7 concludes.

1.1. Related Work

This paper contributes to several literatures.

First, a large body of scholarship documents the barriers that the cost of legal representation creates for low-income individuals (Rhode, 2004; Sandefur, 2019; Hadfield, 2016; Legal Services Corporation, 2022), and experimental evidence shows that legal assistance materially improves outcomes for those who receive it (Greiner et al., 2013). While we do not have demographic data of filers, we provide what is, to our knowledge, the first large-scale quantitative documentation that a general-purpose AI technology has altered the rate at which individuals choose to self-represent in federal court.

Second, we also make a modest methodological contribution to empirical work on civil litigation. In work concerned with case outcomes, economists typically sample historical cohorts where all cases have long since terminated (e.g., Alesina and La Ferrara, 2014; Cohen and Yang, 2019; Ash et al., 2026),⁶ by construction avoiding selection bias from cases that have not yet terminated. Empirical legal studies is more interested in engaging outcomes of recent cases but largely does not appear to account for this termination bias. For example, Lahav and Siegelman (2019), measuring case win rates through 2017, flag the issue in footnote 1 but treat the bias as immaterial given their interest in long-run changes. Levy (2018), on a sample through 2017, similarly acknowledges the concern but simply drops pending cases. The closest paper methodologically is Collinson et al. (2024), which applies a similar observation-window filter, though to labor outcomes after eviction filings rather than to case terminations themselves. Our approach (described in the methods, Section 2.3) is one (conservative) way to deal with the problem from measuring recent case outcomes.

Third, randomized and quasi-experimental studies document productivity gains from LLMs in writing (Noy and Zhang, 2023), customer service (Brynjolfsson et al., 2023), consulting

⁶Related economics work on legal outcomes (Kleinberg et al., 2018; Arnold et al., 2018) further sidesteps the issue by studying outcomes observed at the point of judicial decision (e.g., bail, arrest) rather than intra-case or at termination.

(Dell’Acqua et al., 2023), job search (Wiles et al., 2025), software development (Sarkar, 2025) and entrepreneurship (Otis et al., 2024), with gains often concentrated among lower-ability workers (the exceptions being the last two references). Our setting differs in that the beneficiaries are not necessarily professional workers but individuals accessing a formal institution, and in that we don’t report the outcomes of an experiment with a particular model. Moreover, on the individual’s problem, our results are more consistent with an extensive-margin effect where results in these experiments are often on an intensive-margin.

Finally, an all-important question remains about the capability of models and their downstream effects on the aggregate economy (Acemoglu, 2025; Aghion and Bunel, 2024; Eloundou et al., 2024). On one hand, benchmarks are useful but are potentially weak proxies for improvements in model capabilities (Whitfill et al., 2026). On the other, improving model capabilities seem to carry the risk of real labor displacement (Brynjolfsson et al., 2025), though the heterogeneous effects may differ dramatically across occupations and cohorts (Althoff and Reichardt, 2026; Tucker, 2026).⁷ Our view is that the best indicator of model progress is its ability to do real economic work, and that this ability is precisely the important question to answer if we’re concerned about large labor impacts from LLMs, but that indicators of this work are i) lagging (Brynjolfsson et al., 2021), ii) bottlenecked by data, and iii) difficult to view given partial task automation even with data. This paper documents realized “market-level effects” of LLM diffusion, driven not by the agents of a single firm or actor, but by the uncoordinated actions of many individual people making optimization decisions in their own interest.

2. Data and Methods

2.1. The Federal Civil Litigation Process

Federal civil cases are governed by the Federal Rules of Civil Procedure, which impose uniform pleading, discovery, and motion practice requirements across all 94 U.S. district courts (our data universe). A case begins with the filing of a complaint, which must identify the court’s basis for subject-matter jurisdiction, the legal theory supporting the claim, and the factual basis for relief. Cases are classified at filing by a Nature of Suit (NOS) code, a four-digit

⁷Johnston and Makridis (2026), by contrast, find positive output, employment and wage effects of AI-exposure/adoption, though they ultimately conclude that AI is a capital-biased technology.

taxonomy that distinguishes among approximately 80 substantive legal categories ranging from civil rights and consumer credit to patent infringement and securities fraud.

The *pro se* litigant occupies a formally protected but practically disadvantaged position in this system. Courts are required to construe *pro se* filings liberally and many districts maintain self-help centers offering procedural guidance.⁸ Nevertheless, *pro se* litigants consistently achieve worse case outcomes than represented parties, even after conditioning on case type and court (Greiner et al., 2013; Steinberg, 2023). The gap reflects not only deficiencies in legal knowledge but also disparities in the capacity to navigate procedural requirements, respond to motions, and engage in discovery. Courts are aware of this disparity, and as a result provide *pro se* cases with disproportionate judicial attention from screening complaints for frivolousness and construing deficient pleadings liberally to providing the procedural guidance that counsel would otherwise supply. For this reason, *pro se* litigation has long been understood as a particularly large source of court burden (Wood, 2016).

Finally, it is important to place federal courts in the context of state and municipal courts. Federal courts are the hardest test for potential *pro se* litigants, given the barriers to entry are high along every margin. Financially, the federal filing fee of \$405 is about double the median state-court fee, and only the largest states (e.g., California, Florida) charge comparable amounts. Procedurally, federal courts have higher standards for complaints, so that the procedural barriers to surviving the pleading stage are comparatively higher than state courts, a constraint that often binds on *pro se* litigants the hardest (Eisenberg and Clermont, 2014).⁹ Substantively, federal cases are often high stakes. Federal courts can hear two types of cases, either i) “federal question jurisdiction,” or cases involving federal statutes or actions; or ii) “diversity jurisdiction,” concerning parties of multiple states with more than \$75,000 in dispute. That the *pro se* filing surge we document is occurring in this demanding context suggests the effect in state and local courts—which handle over 90% of civil litigation with simpler rules and high baseline self-representation rates (National Center for State Courts, 2015)—is likely larger still.

⁸Judges must apply the “liberal construction” standard established in *Haines v. Kerner* (1972), in which federal courts try to read *pro se* filings charitably, often overlooking technical defects.

⁹For the interested reader, federal complaints (regardless of whether a plaintiff is self-represented or not) require “factual plausibility,” whereas most state courts accept “notice pleading” (a complaint just has to tell the other side what the dispute is about). This is also why judge burden is higher from the start of the case in a federal case, where a judge may have to adjudicate whether a complaint is factually plausible in the first place.

2.2. Data Sources

2.2.1. FJC Integrated Database (*Duration and Disposition Analysis*)

For the analysis of trends in case filing, case duration, and case outcomes, we use the Federal Judicial Center Integrated Database (IDB), which records case-level information on all civil cases terminated in all U.S. district courts (Federal Judicial Center, 2025). Our data spans data files from FY2005-2025, which we combine into a single deduplicated dataset indexed by unique case identifiers.

The resulting sample contains approximately 3.7 million cases with valid filing dates, termination dates, and positive recorded durations. For each case we observe the NOS code, the district, the filing date, the termination date, the representation status (PROSE field, coded 0 for represented, 1 for plaintiff *pro se*, 2 for defendant *pro se*, 3 for both parties *pro se*), the disposition code (DISP), and the judgment outcome (JUDGMENT).

2.2.2. AOUSC C-13 Table

For filings data, we augment the IDB data set with the C-13 table published by the Administrative Office of U.S. Courts, which records civil case filings by district, year, and representation status. This dataset includes only district filing data aggregated by year. While both limited as only annual, district-level aggregated data, this data source primarily acts as a source of robustness for our results in Section 3.1.

In both the IDB and C-13 datasets, we exclude two populations of case filers. First, we exclude cases from prisoner filings (cases arising under 28 U.S.C. §§ 2241, 2254, 2255, and related statutes), which constitute a qualitatively distinct population responding to different institutional incentives and under different *pro se* requirements. We also exclude organized mass-litigation events.¹⁰ We do so as such organized mass-filing events (often related to product recalls) depart from the qualitative features that characterize typical civil *pro se* litigation, thereby obfuscating from the interesting economic and social phenomena at play.

¹⁰In the IDB set, this exclusion is cases flagged as multidistrict litigation transfers (MDLDOCK $\neq -8$) and MDL-related origin codes (ORIGIN $\in \{6, 13\}$), which capture product-liability mass torts consolidated in a single district. In C-13, we exclude the FL,N district-year cells, which exhibit a mass-filing anomaly of 196,000 *pro se* filings in FY2020 (over 1,400 times the district’s baseline).

2.2.3. PACER docket entries

We also access metadata for 46 million PACER docket entries and match them to IDB cases to characterize docket activity within each case. Federal district courts differ in how they publish their docket data. For the 54 district courts (of 94 total) that provide full PACER feed coverage, we achieve a 99.8% match rate to the underlying case data. Despite attrition in the sample due to losing the 40 districts that do not provide “full-feed” data, our match is comprehensive for approximately 80% of all non-prisoner *pro se* filings.

There is a small gap in our 2024 Q4 data localized in represented cases. This does not affect our ability to observe cases as a whole (filings, terminations, outcomes are all left unchanged) but affects the intra-case data (namely, number of entries within a case, as discussed in Section 4.2).

2.3. Methods

2.3.1. Correcting for right-censoring and termination bias

Our analysis concerns patterns associated with both case filings and case outcomes. However, almost tautologically, we can only observe outcomes for cases that have terminated. So, any analysis of outcomes must account for selection from right-censoring. To illustrate, imagine there are only two equally likely cases, cases that resolve in 6 months c_A and cases that resolve in 18 months c_B . Suppose further that we can observe data only through December 31, 2025. For cases filed sufficiently far in the past, we would observe that cases have an average duration of ≈ 1 year (including both c_A and c_B terminations), whereas as we approach the right-edge of the time series we would observe an average duration of ≈ 6 months (including only c_A terminations), and thereby incorrectly assume that recent cases resolve doubly fast! Indeed, if we compute any outcome measure—say, number of docket entries per case—on the right-edge of the sample and compare it to the left-edge to view the difference in that measure over time, we’ll be polluting that comparison by also implicitly measuring the difference between the number of entries on average in $\{c_A\}$ and $\{c_A, c_B\}$ cases. This is the termination-bias problem.

We implement a horizon-specific, nonparametric estimator to account for right-censoring of outcomes. Let R_i denote case i ’s available follow-up time (December 31, 2025 minus filing

date) and let T_i denote its duration (termination date minus filing date). Fix an observation window W and let $t \in [0, W]$ denote the evaluation horizon within that window. We restrict attention to cases with $R_i \geq W$, so that every retained case has had the same opportunity to resolve within W days. For duration analysis, this results in a completion function

$$\hat{G}(t; W) = \frac{\sum_i \mathbb{1}\{R_i \geq W, T_i \leq t\}}{\sum_i \mathbb{1}\{R_i \geq W\}}, \quad 0 \leq t \leq W.$$

Returning to the vignette, setting $W = 180$ allows us to compare only filing cohorts whose cases have all had at least 6 months to resolve (i.e., only comparing c_A cases across time). For other outcomes (e.g., disposition or judgment analyses), we apply the same $R_i \geq W$ restriction and compute outcome shares among the subset of retained cases with $T_i \leq W$.

Choosing W imposes a simple tradeoff. Smaller values of W retain filing cohorts closer to the right edge of the data set, but only identify outcomes among fast-resolving cases. Larger values of W capture slower-resolving cases, but exclude more recent filings. In our setting, this tradeoff is especially salient because large values of W mechanically remove the most recent post-AI cohorts—if tech diffusion is less than immediate, these are the cohorts where changes in outcomes from AI should be most visible. Through the main body (Section 4) we report outcomes along the 6-month horizon ($W = 180$ days) and report robustness across the full suite of windows $W \in \{90, 270, 365\}$ in the Appendix.

3. Who is filing, and in what cases?

We document a steep rise in the number and share of *pro se* filings in federal civil court since FY2022, driven primarily by a growing number of self-represented plaintiffs. We also find that this rise concentrates in types of cases where document production constitutes the primary cost of filing, and that this rise has taken place almost uniformly across states.

3.1. The national *pro se* share and count

Figure 1 illustrates this rise using both of our primary datasets. Using the C-13 database, the *pro se* share of non-prisoner filings hovers around 11% from FY2005 through FY2022, a level Gough and Taylor Poppe (2020) describe as “remarkably unremarkable.” In FY2025 it

reaches 16.8%. The change is continuous and steep rather than a one-year anomaly: shares rise from 11.6% in FY2022 to 11.4% in FY2023 (flat), 14.5% in FY2024, and 16.8% in FY2025. Represented filings do not fall. These contemporaneous trends can be interpreted in two ways:

1. First, it is possible that the additional *pro se* cases are new entrants to the federal courts who, without LLMs, would never have entered.
2. Alternatively, it is possible that the advent of capable AI tools induces parties to substitute away from professional legal representation and to file their lawsuits *pro se*. However, this substitution effect must have been dominated by the entry of new represented parties, because we do not observe a net decrease in the number of represented cases (Figure 2).

Pro se share of federal civil litigation, deviation from pre-AI mean

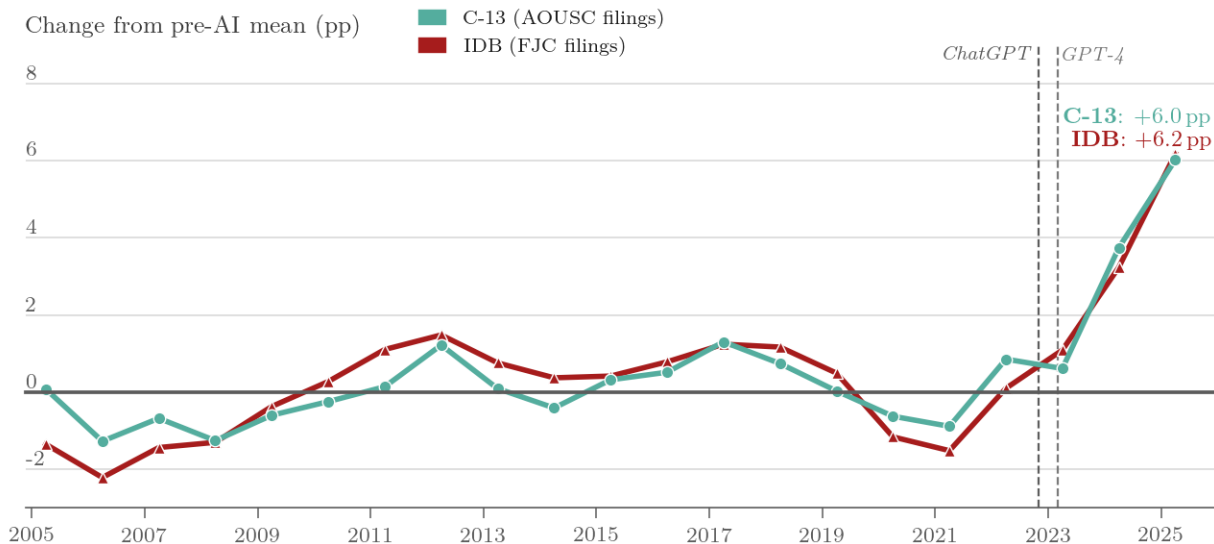


Figure 1. National non-prisoner *pro se* share, FY2005–FY2025. Twenty years of stability at an approximately 11% steady-state followed by a sharp post-2022 rise. Source: AOUSC C-13 and FJC IDB.

Figure 2 plots the number of represented and *pro se* cases filed by fiscal year in federal district courts. Two facts stand out.

First, represented filings are essentially constant over this period. Second, *pro se* filing counts have broken out of their long-run range. *Pro se* filings counts average 23,210 cases per year from FY2005 to FY2022. Beginning in FY2023, *pro se* filings jump: 27,370 in FY2023, 31,478

Federal civil complaint filings, by representation status

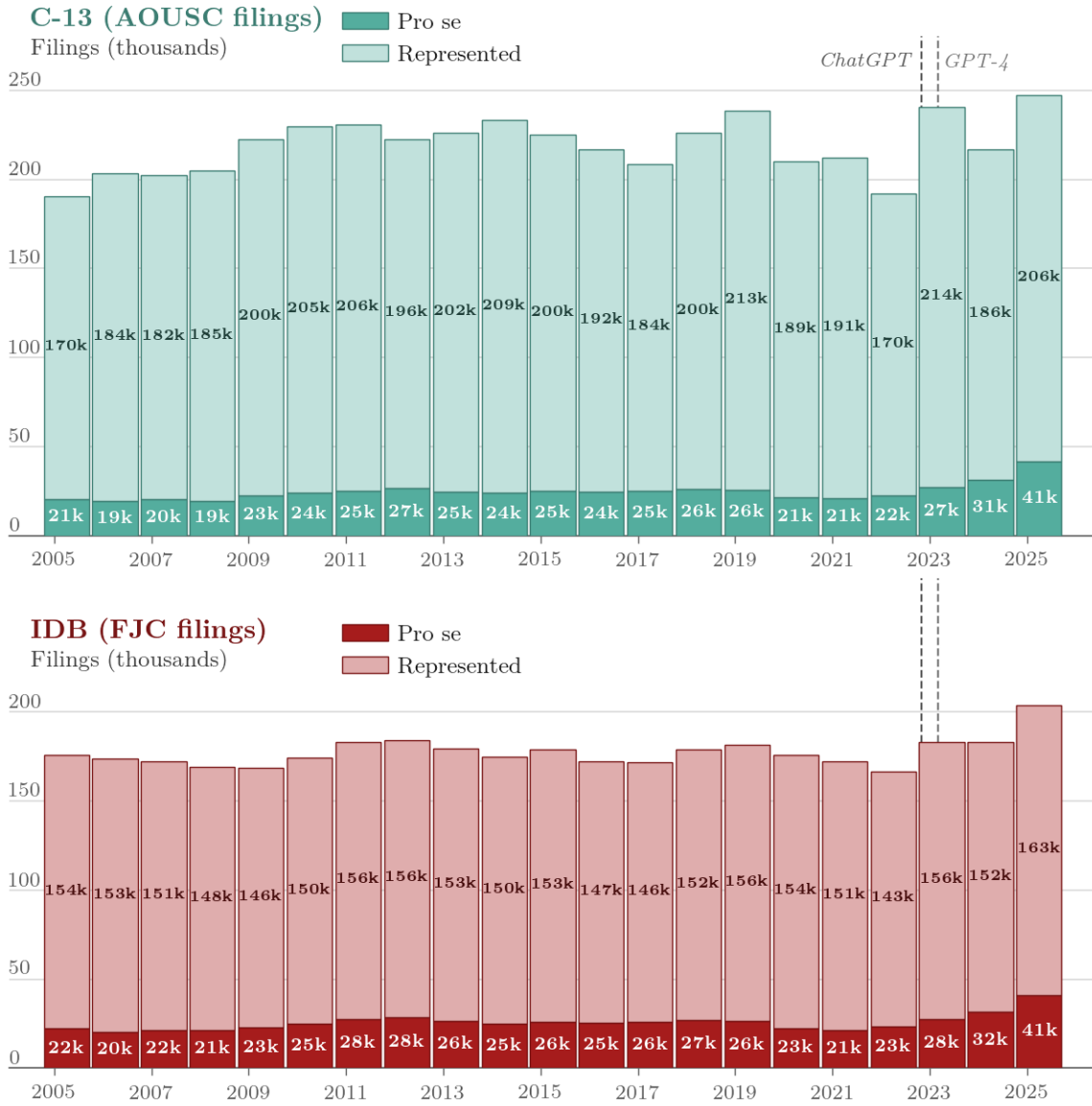


Figure 2. Non-prisoner federal civil filings, FY2005–FY2025. Represented filings are essentially flat; *pro se* filings step up sharply starting in FY2023. Source: AOUSC C-13 and FJC IDB.

in FY2024, and 41,490 in FY2025. The FY2025 count is almost double the pre-AI mean.

Using the AOUSC filing data, by FY2025, total civil litigation had increased by roughly 31,170 cases relative to the pre-AI mean (a 14.4% increase in counts). *Pro se* filings alone accounted for 18,280 of those additional cases. Despite representing a share of all filings an order of magnitude smaller than represented cases, new *pro se* filings account for 59% of

the growth in civil filings. We interpret this as being consistent with the conjecture that the widespread adoption of generative AI tools is driving increased filings by those who gain disproportionately from the new technology at the filing stage: non-lawyers.

3.2. Plaintiffs, not defendants

A *pro se* case can arise for two very different reasons. A plaintiff may choose to file without counsel because they have a claim they want to pursue and cannot or will not incur the costs of hiring a lawyer to pursue it. A defendant may appear in court without counsel because they were sued and cannot afford a professional defense. Because plaintiffs initiate civil legal proceedings, their choice to forgo professional representation is materially different than defendants' choice to do the same. Plaintiffs have the null outside option (not entering), whereas defendants, faced with a complaint, have the outside option only of being represented. Figure 3 decomposes the annual *pro se* case count into these roles using the IDB PROSE field, which distinguishes plaintiff-side, defendant-side, and both-sides self-representation.

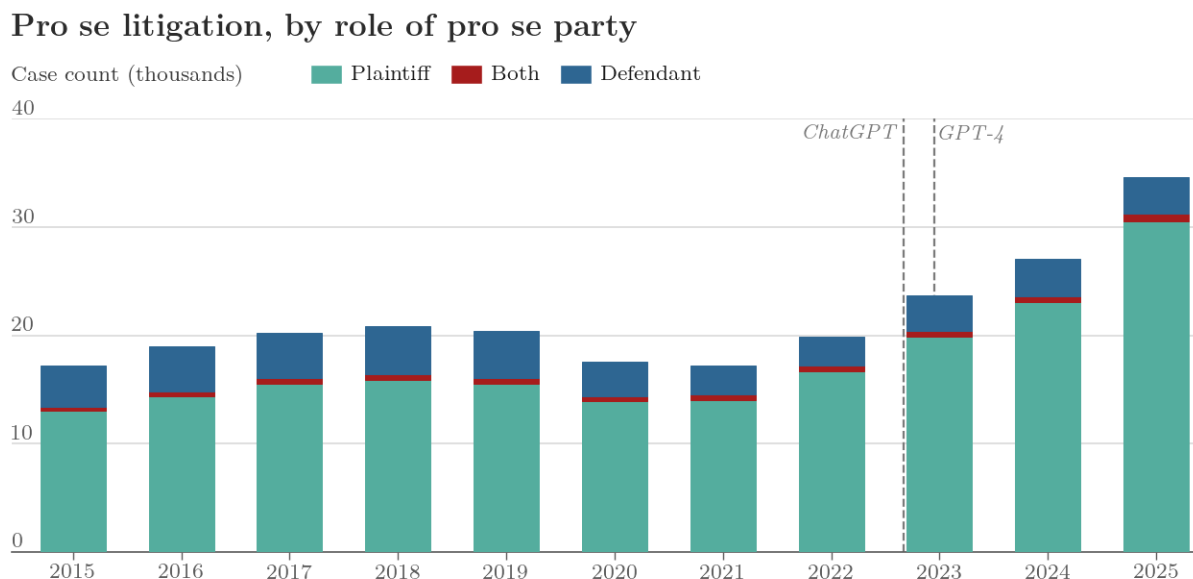


Figure 3. Pro se case counts by role, FY2015–FY2025. The post-AI rise loads almost entirely on plaintiff-side *pro se* filings. Source: FJC IDB, non-prisoner non-MDL.

The post-AI rise in self-representation is a plaintiff story. Plaintiff-side *pro se* case counts averaged 19,705 per year from FY2015 to FY2022 and reach 39,167 in FY2025, nearly doubling. Defendant-side *pro se* counts fall slightly over the same window, from 4,650 to

3,896. Both-sides *pro se* cases grow modestly, from 775 to 1,081. The increase in the total *pro se* count since FY2022 is therefore accounted for by plaintiff-side entry.

3.3. Simple cases, not complex ones

If generative AI has lowered the cost of producing filing documents, the rise in *pro se* entry should concentrate in case types where filing-stage document production represents the dominant cost in initiating and pursuing a lawsuit. For civil rights or consumer credit cases, which often resolve at or shortly after the pleading stage, AI can complete at *de minimis* cost many of the tasks that a lawyer would complete for their client, but are extremely demanding for even educated laymen. For patent infringement or securities fraud, which require sustained discovery, adversarial motion practice, and expert testimony, tasks are less separable, and so remain more immune to task-wise automation from *pro se* litigants. So, does the post-AI rise in *pro se* litigation actually load on the case types where filing-stage help matters most?

3.3.1. A brief model of simple and complex cases

Before turning to the data, we exposit one way to think about simple versus complex cases and how AI may operate as a shock to filing. We adapt a simple model of home production from Becker (1965).

Consider a potential litigant i with a claim in case type n . Normalize the outside option (not filing) to 0. In the market, the litigant has the ability to pay representation p_n per unit of legal work and receives expected case payoff $\mathbb{E}[Y_{in}^R]$:

$$U_{in}^R = \mathbb{E}[Y_{in}^R] - p_n.$$

The litigant also has the ability to produce legal services with their own resources and wherewithal at effort level e . The cost to the litigant of supplying this effort is increasing and convex in e . The expected case payoff $\mathbb{E}[Y_{in}^P(e)]$ is increasing and concave in e . The litigant then chooses e to maximize:

$$U_{in}^P(e) = \mathbb{E}[Y_{in}^P(e)] - c_{in}(e),$$

yielding the optimal effort e_{in}^* and resulting utility $U_{in}^{P,*}(e)$.

The litigant files if $\max\{U_{in}^{P,*}, U_{in}^R\} \geq 0$, and conditional on filing chooses representation if $U_{in}^R \geq U_{in}^{P,*}$. Equivalently, the litigant retains counsel iff $p_n \leq V_{in} = \mathbb{E}[Y_{in}^R - Y_{in}^P(e^*)] + c_{in}(e^*)$, where V_{in} is the litigant's private value of representation (the outcome premium from hiring a lawyer plus the personal effort the lawyer saves the litigant). The pre-AI *pro se* share in case type n is then the fraction of (filed) litigants for whom the lawyer is not worth the price:

$$PS_n^{pre} = \Pr(V_{in} < p_n \mid \text{filed}, n) = F_n(p_n),$$

where F_n is the case-type-specific CDF of V .

We model AI as a productivity shock to the home production of legal documents. More precisely, this is that AI lowers the shadow cost of self-supplied legal-document production, so that the marginal cost $c'_{in}(e)$ falls and the marginal product $Y_{in}^{P'}(e)$ rises at any given e . Both effects reduce the litigant's private value of representation V_{in} . Call this productivity benefit a , so that $V_{in}^{post} = V_{in}^{pre} - a$. Then, the set of litigants that enter are those just barely on the margin pre-AI, $V_{in}^{pre} \in (p_n, p_n + a]$, and the change in *pro se* share in n is

$$\Delta PS_n = F_n(p_n + a) - F_n(p_n) \approx a \cdot f_n(p_n).$$

This gives the interpretation of the pre-AI *pro se* share we use in the rest of this section. It is not a literal measure of doctrinal complexity. It is a revealed-preference measure of how often, in a given case type, the private value of representation fell below its price. High *pro se*-share categories are categories in which self-help was already relatively attractive before AI, whether because the gain from counsel was small or because the effective price of counsel was large. Low *pro se*-share categories are those in which representation was more worth purchasing. In the language of the model, the CDF $F_{\text{simple}}(\cdot)$ is first order stochastically dominated by $F_{\text{complex}}(\cdot)$ (at least over the region near p_n).

We have two additional notes.

First, note that this toy model normalizes lawyer-side production so that we only have to consider AI improvements on the *pro se* litigant's side. If we impose the same features on the

lawyer’s production function of legal documents—so that $Y_{in}^R(e_\ell)$ is concave in e_ℓ and that $c_{in}(e_\ell)$ is convex—then the same dynamics play out. Because lawyers are already producing adequate documents in the absence of AI tools (exerting greater effort e_ℓ), the marginal net gains in the difference $Y(\cdot) - c(\cdot)$ is smaller than the household’s net gain when subject to the same AI shock.

Note further that LLMs may provide a second benefit (or cost), namely that they change a person’s information environment. A potential litigant who pre-AI did not know they had a cognizable claim now has access to a LLM that can advise them, and may file a claim they would not have known to file before. We load all this in for now into $Y(\cdot)$ but future work may want to unpack this more carefully. This model admits an increase in *pro se* cases being filed post-AI by both “entrants” new litigants entering the legal system and “switchers,” previously represented litigants choosing to file *pro se*.

3.4. Simple cases rise the most in the post-AI period

We use the pre-AI *pro se* share of each Nature of Suit (NOS) category as a revealed-preference measure of how amenable the case type is to self-representation. We separate NOS categories purely along their pre-AI *pro se* share.

NOS categories whose pre-2022 *pro se* shares were high are categories where many litigants, even before GPT-4, had concluded that self-representation was feasible.¹¹ NOS categories whose pre-2022 *pro se* shares were low are categories where self-representation was rare, presumably because the case type required more than filing-stage work. We split the 86 non-prisoner NOS codes at the case-weighted median of pre-AI *pro se* share, yielding 48 “simple” and 38 “complex” categories. The simple category contains, among others, civil rights (NOS 440, 51% pre-AI *pro se*), foreclosure (220, 31%), and employment discrimination (442, 18%). The complex category contains patent (830, 3%), product liability (365, 4%), and insurance (110, 5%). We report the list of NOS codes (with pre-AI average case volume and *pro se* share) in Appendix B.1.

Figure 4 plots the *pro se* share for each category by fiscal year. Additionally, while we split NOS codes into “simple” and “complex” along the case-weighted median of the pre-

¹¹This may have been due to *ex ante* variability in the templatability of certain types of suits, the availability of specialized legal-aid or self-help clinics, and other informational resources.

Pro se share of case-type filings, by simple vs. complex NOS case-type

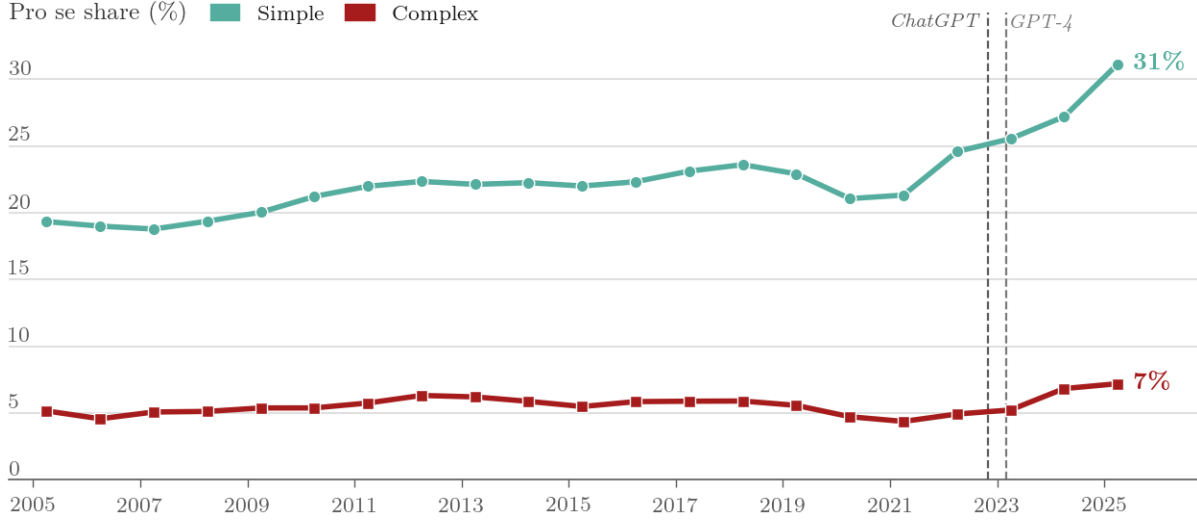


Figure 4. *Pro se* share by fiscal year, split at the pre-AI case-weighted median of *pro se* share. The simple category sees the *pro se* share rise sharply post-2022; the complex NOS *pro se* share barely moves. Source: FJC IDB, non-prisoner, non-MDL cases.

AI *pro se* share of each NOS code, we also present versions of this plot for splits at the {80, 70, 60, 40, 30}-th percentiles as well in Appendix B.3, and see similar separations across these splits.

	Simple	Complex
Federal question	59.6%	47.4%
Original complaint	94.2%	69.8%
Standardized complaint	62.8%	0.0%
Specialist complaint	1.3%	10.9%
Other original complaint	30.1%	58.9%
Review / petition	5.8%	30.2%

Table 1. Structural characterization of the simple/complex split, pre-AI. Each entry is the FY2005–FY2022 share of cases in the simple or complex partition falling in the indicated bucket. “Original complaint” and “Review / petition” partition the universe; the indented rows partition the original-complaint universe. See Appendix B.2 for exact NOS-to-bucket mappings.

What NOS-characteristics distinguish these two groups from each other beyond the *pro se* share used to construct them? Table 1 decomposes each partition along three dimensions derived from the IDB’s jurisdictional and case-type fields. As discussed in Section 1.1, there

are two types of cases in federal courts, and the first row reports federal question jurisdiction, or cases which arise under a federal statute. One thought is that this structure may make it easier for *pro se* claimants to file under. Through the pre-period sample, simple cases are somewhat more likely to be federal question cases (60% versus 47%), but the overlap is substantial—this is primarily because there is substantial heterogeneity in how complicated or burdensome the requirements of different federal statutes are.

More informative is the distinction between original complaints and review or petition matters. An original complaint initiates a new dispute: the plaintiff drafts a factual narrative, identifies a legal theory, and asks the court for relief. A review or petition, by contrast, asks a federal judge to review a decision already made by another body—e.g., from an administrative agency, a bankruptcy court, or an immigration officer—so that the plaintiff’s task in a review case is not to construct a narrative but to identify legal error in an existing record. Nearly all simple-side cases are original complaints (94%), whereas 30% of complex-side cases are review or petition matters—cases like Social Security disability review, ERISA benefit denials, and bankruptcy appeals.

Disaggregating the original-complaint cases further, the sharpest separation comes from whether these cases are “standardized” or “specialized.” The AOUSC publishes national *pro se* complaint templates,¹² and we define “standardized” original complaints as the NOS categories that correspond to these templateable cases. In contrast, we define the “specialist” category as case types where filing itself requires credentials or procedural knowledge beyond ordinary pleading: a statutory pleading standard stricter than plausibility (securities under the PSLRA), a registration prerequisite (copyright), a mandatory seal period (False Claims Act), or even a separate federal bar altogether (patent). Both definitions are intentionally conservative: on the simple side, consumer credit, fraud, and telephone consumer protection cases are highly templatable in practice but lack a corresponding AO form, so they remain in the residual “other original complaint” category. Further details on the classification, including exact NOS-to-type mappings, are provided in Appendix B.2.

Even with these conservative definitions, we see that simple filings are more likely to be in templateable, procedural cases where good filing constitutes a larger share of the case’s total required work; complex cases are more likely to be in diverse, non-templateable cases that

¹²<https://www.uscourts.gov/forms-rules/forms/civil-pro-se-forms>.

may require sustained expertise. We see this as a ratification of our pre-AI *pro se* share split as capturing some notion of how LLMs should be expected to affect filing shares upon the diffusion of LLMs.

3.5. Geographic spread

The third question arises about *where* these cases are being filed. If the *pro se* rise reflects something local—a single district’s policy change, a courthouse that became especially welcoming to self-represented filers, or a state bar association or an area law school’s initiative—we would expect the increase to be concentrated in a few places. If instead it reflects a technology shock available everywhere at the same time, we would expect the increase to show up almost everywhere at once.

Change in *pro se* filings to FY2025 vs. pre-AI peak

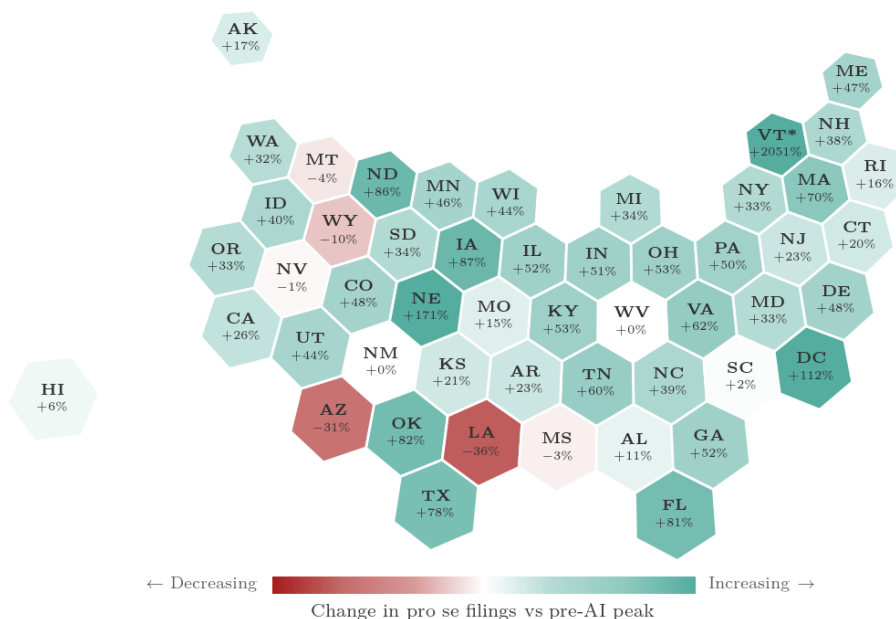


Figure 5. FY2025 non-prisoner *pro se* filing count relative to each state’s peak pre-AI year (highest annual count across FY2005–FY2022). Darker cells indicate states where FY2025 exceeds the pre-AI peak. A case study on Vermont is provided in Appendix A. Source: FJC IDB.

Figure 5 aggregates district-level *pro se* filing counts to the state level (there are no federal district courts that cross state boundaries) and plots FY2025 relative to each state’s peak pre-AI year. For a state to appear “teal” in this figure, FY2025 *pro se* filings must exceed the single highest pre-AI year in that state’s own history. All but six states clear this maximally

rigorous bar.^{13,14} The rise in *pro se* cases is not driven by a handful of outlier districts. It is geographically pervasive in a way that is difficult to reconcile with any explanation local to a specific jurisdiction.

4. What are courts doing with these cases?

The previous section documented a large and geographically pervasive increase in federal *pro se* civil filings, concentrated especially in simple case types. But what happens to those cases once they enter the federal system? If AI is helping self-represented plaintiffs draft complaints but nothing else, the additional cases should resolve at around the same pace as pre-AI cases and leave little trace on the judicial system. If AI is helping *pro se* parties sustain litigation further into the process, durations should shift, dispositions should change, and motion activity should rise.

Three caveats apply throughout this section. First, all of our outcome measures (within-case activity, case durations, dispositions, and judgments) are subject to right-censoring: we observe outcomes only for cases that have terminated as of December 31, 2025, so cases that would take a long time to resolve are under-represented in the most recent cohorts. To correct for this, we apply observation-window corrections throughout, following the method described in Section 2.3. Second, the COVID-19 pandemic is a large confounder for any outcome measured on filing cohorts between March 2020 and late 2021; we read the post-AI series against the pre-COVID baseline (FY2015–FY2019) and flag COVID artifacts we assess them to have material implications for our analysis. Third, the PACER docket data used for intra-case activity measures covers the 54 federal districts with full-feed PACER access, a subset that accounts for roughly 80% of non-prisoner federal civil cases. Within these full-feed districts, we have a 99.8% match rate with the PACER docket data, and we report outcomes from these full-feed districts.

¹³Large states often contain multiple federal judicial districts, so we also report this summary at the district level in Appendix B.4. In sum, only 14 of 94 judicial districts (15%) experience a contraction in the *pro se* filing count from their pre-AI peaks.

¹⁴A notable positive outlier is Vermont, which also reflects strategic “forum shopping” by a specific group of *pro se* immigration litigants rather than a broad trend. We find Vermont to be a particularly interesting district, and provide a more detailed case study in Appendix A. Excluding Vermont leaves every result in the paper qualitatively unchanged, and we also report plots to this effect.

4.1. Case durations are unchanged

The simplest outcome to measure is how long *pro se* cases take to resolve. If the new *pro se* filings are meritless and being dismissed on the papers, the share of cases that are resolved in an observation window should increase. If they are sustaining litigation further into the process or judges are overburdened, the share of cases resolved in an observation window should decrease. Figure 6 plots the share of *pro se* cases resolved by 180 days ($W = 180$) from filing, computed at the monthly filing cohort level.

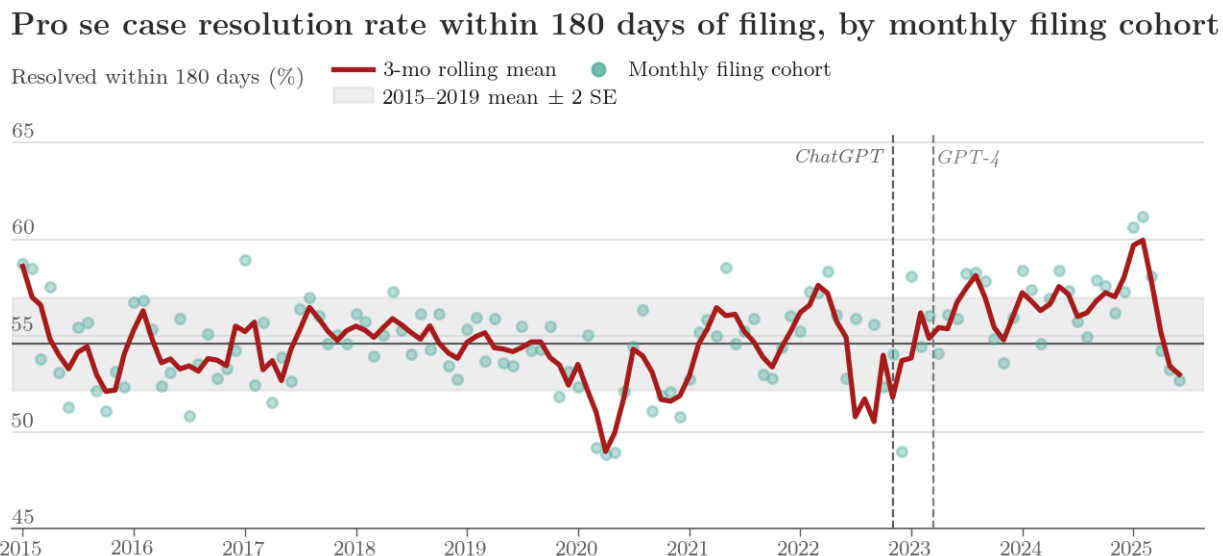


Figure 6. Monthly cohort *pro se* case resolution rate at 180 days from filing, FY2015–FY2025. Each dot is one filing-month cohort. The post-2022 series sits within the pre-COVID noise band. Dashed lines mark ChatGPT (Nov 2022) and GPT-4 (Mar 2023). Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

The post-AI *pro se* resolution rate sits essentially within the pre-2020 noise band.¹⁵ Cases filed in FY2023–FY2025 resolve no faster or slower than cases filed in FY2015–FY2019.¹⁶ The COVID period produces a transient dip in resolution rates that persists through roughly mid-2021 and recovers afterward. We observe that the post-AI period case resolution rate marks a return to the long-run baseline. By the simplest measure of court throughput—the

¹⁵Figure 6 presents the 180-day case resolution rate for *pro se* cases. Appendix B.5 presents the 90- and 365-day horizons, as well as companion plots for represented cases. All show the same general trend, of case durations being largely unchanged in the post-AI period.

¹⁶We use the FY2015–FY2019 window here, rather than the usual FY2015–FY2023 pre-AI window because COVID appears to have a meaningful effect on durations (whereas its effect on other measures of interest is less marked) over all observation windows.

time from filing to termination—AI has made no visible mark.

4.2. But docket activity inside cases has risen sharply

Case duration statistics may generally summarize the arc of cases but say nothing about what happens along the way. A second class of measures looks *inside* cases: the docket entries generated as a case progresses—filings, orders, responses, scheduling events. These entries are the measures that capture the burden a case places on the court while it runs its course, and they are the indicia upon which AI has left the clearest mark.

Figure 7 plots the number of docket entries generated in the first 180 days after a case is filed, quarterly, separately for *pro se* and represented cases. The left panel plots the total number of entries per court; the right panel presents the number of entries per case. Results at the 90-day horizon are qualitatively similar and are reported in Appendix B.6.

By 2025Q2 (the most recent fully observable quarter at the 180-day horizon), the number of entries per court from *pro se* cases is up 158% relative to the pre-AI mean. These effects don’t begin until 2025 (likely because technology diffusion takes time), so that averaging over the post-period understates these results. The left column includes two effects, both that the number of cases per court is increasing (our headline result) and that the number of entries per case is increasing (an intensity result). The right column, plotting the entries per case in the first 180 days of a case’s life, makes this explicit. By 2025Q2, *pro se* entries per case reach 23.3, a 38% rise compared to the pre-AI mean of 16.9.

4.2.1. Symmetric, not one-sided

These results are smaller in absolute magnitude but qualitatively similar for represented cases. Represented cases also generate more docket entries per case than they did before: 22.5 in 2025Q2 versus 18.2 pre-AI, a 23% rise in the first 180 days. This can be explained by attorneys using LLMs too, suggesting LLMs are having an important intensive-margin effect for both groups as a general technology that lowers drafting costs.¹⁷ It’s also important to

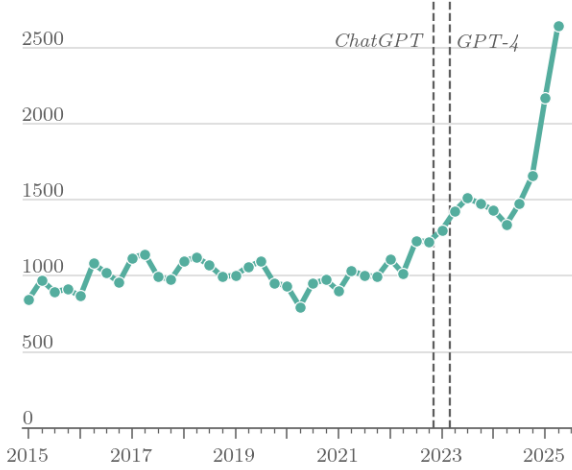
¹⁷Note that the legal profession has strong norms (and sanctions) that might discourage the adoption of consumer-grade publicly available tools like ChatGPT by attorneys. However, we use the release of ChatGPT and GPT-4 as proxies for the advent of powerful AI tools that have spawned a slate of “legal-tech” startups marketed to lawyers and law firms. Even if these tools are not being used to wholesale draft filings, by reducing the cost of other tasks of legal practice they may lead lawyers to reallocate more time or effort to producing written submissions.

Civil litigation docket-entry burden, by representation status

Pro se

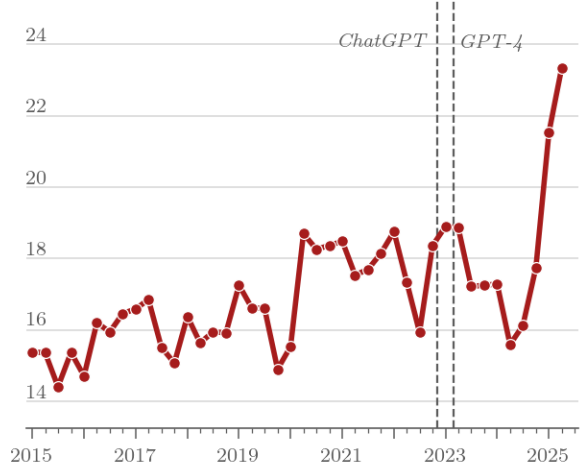
180-day entries per court

Entries per court



180-day entries per case

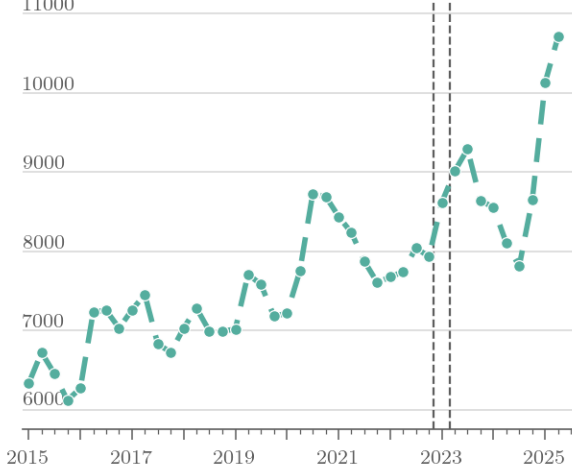
Entries per case



Represented

180-day entries per court

Entries per court



180-day entries per case

Entries per case

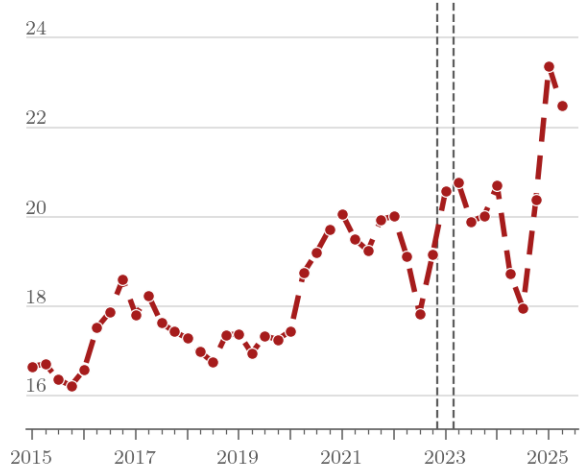


Figure 7. Docket entries in the first 180 days of a case, quarterly, for 54 full-feed districts. Left: entries per court; right: entries per case. Top row: *pro se*; bottom row: represented. Dips in 2024 Q4 can be attributed to a data gap particularly affecting represented cases, as discussed in Section 2.2.3. This gap only affects metadata on intra-case activity (e.g., entries within a case, not outcome or filing data). Source: PACER docket metadata; FJC IDB.

note that we have a data issue especially localized in represented cases for 2024Q4 (discussed also in Section 2), where our ability to view intra-case data (namely, case entries and entry types) is limited. This data issue is repaired by 2025Q1, and we are working to obtain additional sources of data to repair it for 2024Q4.

The notable pattern is not that one side is reacting to the other. It is that both sides are generating more docket activity at once, so the total burden a case places on the court during its first 180 days is rising on both sides of the bar, not just on the *pro se* side of it. This pattern is also consistent with the notion that LLMs are a widely available general purpose technology. Results are similar at various observation-window horizons and reported in Appendix B.6.

4.2.2. Concentrated in simple case types

Furthermore, if AI is the driver of these changes, the docket-activity rise should also load on the same case types where the *pro se* filing rise loaded—the simple NOS categories. Figure B15 repeats the per-case and per-court analysis separately for the simple and complex categories, using the same pre-AI *pro se*-share partition as described in Section 3.3.

The concentration is sharp and consistent with our other findings. In simple NOS cases, *pro se* entries per case rose from 16.4 pre-AI (2017Q1–2022Q3) to 24.3 in 2025Q2 (a 48% increase), and total *pro se* entries per court rose from 802 to 2,234 per quarter in the first 180 days (+179%). In complex NOS cases, entries per case rose only modestly from 19.0 to 19.2 by 2025Q2 (essentially flat), while entries per court rose from 227 to 413 (+82%)—growth coming almost entirely from the higher volume of complex-side filings rather than per-case intensification. These patterns substantiate the lowered cost of *pro se* litigation loading into simpler suits, where AI-augmented parties can more readily produce written submissions early in a case’s lifespan. The additional docket activity documented in the previous subsections is almost entirely a simple-NOS phenomenon. It tracks the same partition as the *pro se* filing rise, which is consistent with a common underlying cause. This is also demonstrated starkly in deviations-from-pre-AI-mean terms in Figure B16.

4.3. Case outcomes are unchanged

Coming to the end of a case’s life cycle in the trial court, the last question is of judgment. A natural prediction is that lowering the cost of filing should draw in marginal cases of lower average quality. Holding constant the facts of a prospective lawsuit, if AI makes it cheaper to file a complaint, the complaints that would not have been filed absent AI are weaker than those that were filed. However, we find little evidence that the quality distribution has shifted markedly: post-AI *pro se* cases resolve in approximately the same distribution as their pre-AI predecessors.

Figure 8 plots the quarterly count of *pro se* case outcomes within 180-day filing cohorts, classified into four mutually exclusive categories: *pro se* unambiguously wins, *pro se* unambiguously loses, negotiated/settled, and dismissed by judge. Counts rise across every category (consistent with Figure 2). Quarterly cases resolved within 180 days roughly double from a pre-AI mean of approximately 2,000 to over 4,000 by 2025Q2. Settlements rise from roughly 390 to over 900 per quarter; judicial dismissals rise from roughly 1,200 to over 2,500.

Pro se case outcomes within 180 days of filing

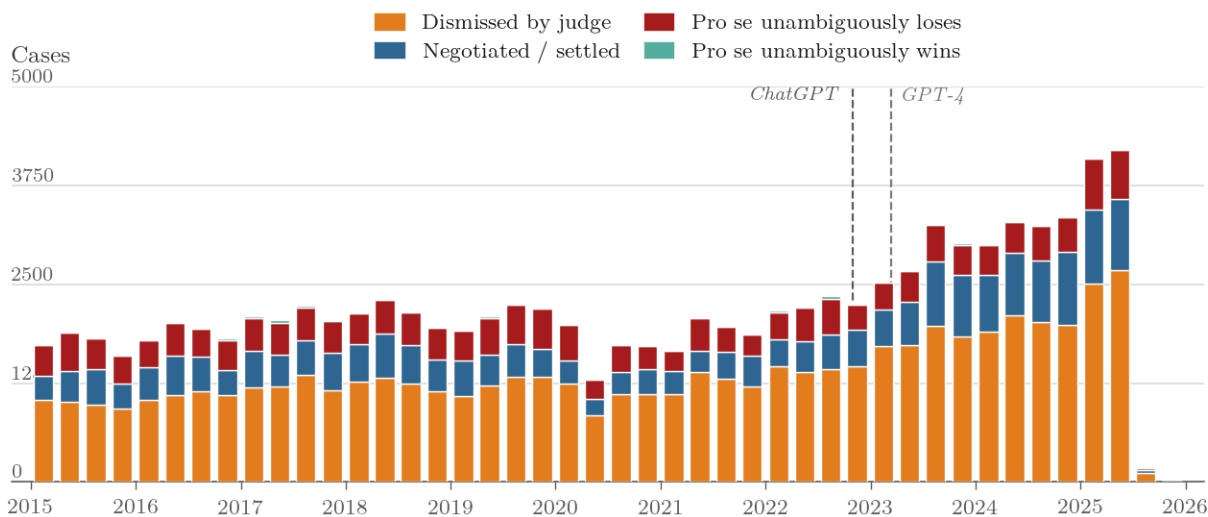


Figure 8. Quarterly count of *pro se* case outcomes within 180-day filing cohorts, by outcome category. Each bar represents cases filed in that quarter that terminated within 180 days. The sharp drop in 2025Q3–Q4 reflects right-censoring correction: most cases filed after approximately July 2025 have not yet had 180 days to resolve. Source: FJC IDB

Figure 9 presents the same data as shares of classified terminations. The composition of

outcomes appears broadly stable. Judicial dismissals account for approximately 60% of resolved *pro se* cases in the pre-AI period and 63% in the post-AI period. Settlements account for roughly 20% pre-AI and 23% post-AI. *Pro se* losses fall modestly from 19% to 13%, and *pro se* wins remain rare throughout (under 1%). The distribution of case quality, as revealed by how cases resolve, has not detectably changed. In contrast to other case characteristics (adversarialism as proxied for by motion activity and case duration), there does not appear to be a clear or distinct departure in case quality. The new, plausibly AI-affected cases are resolving in approximately the same proportions as the old, unaffected ones.

Pro se case outcomes within 180 days of filing

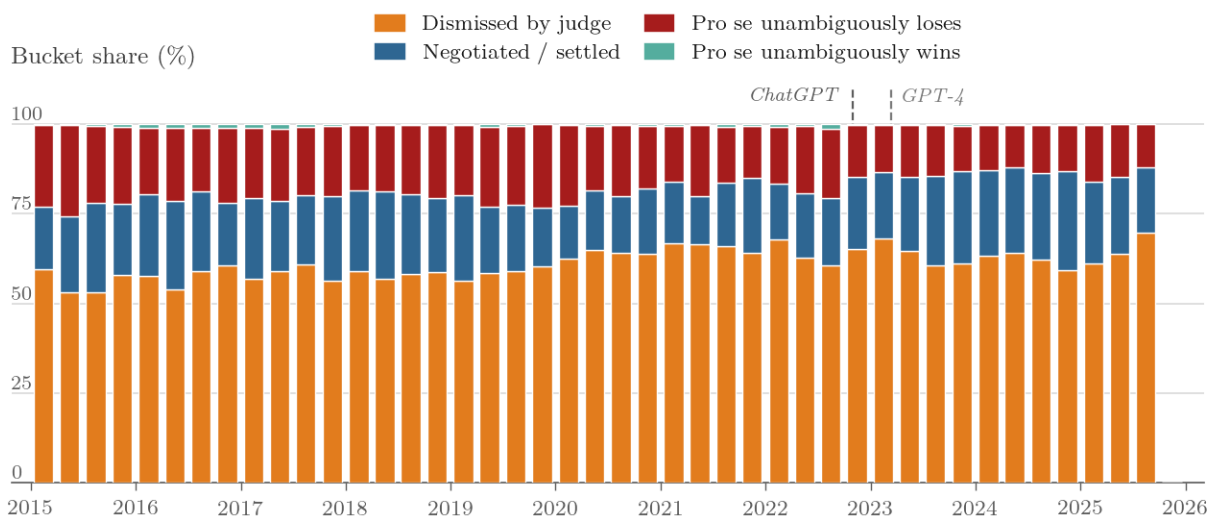


Figure 9. Share of classified *pro se* terminations by outcome category within 180-day filing cohorts. Shares are normalized to sum to 100% within each quarter. The observation-window estimator restricts to cases with at least 180 days of observation, so the rightmost quarters mechanically lose cases that have not yet had time to resolve. Source FJC IDB.

5. Is AI-generated text showing up in the documents themselves?

The results so far are consistent with AI being used by self-represented litigants, but they fall short of being dispositive. The filing increase, the concentration in simple case types, the geographic spread, and the rise in motion activity are patterns we would expect if AI had lowered the cost of legal document production. They could also, in principle, reflect some other shock: a demographic shift, a legal-aid funding change, or a cultural moment we have not considered. To progressively rule out these non-AI alternative explanations, we pull a

random sample of complaint documents from federal courts and directly test whether they are generated by LLM tools.

5.1. Text detector

We use Pangram Labs, a commercial AI-text detection service, for all AI-text detection in this paper. Imas and Jabarian (2025) evaluate four leading AI-text detectors—Pangram, OriginalityAI, GPTZero, and RoBERTa—and find that Pangram is the only detector to satisfy a strict false-positive cap ($\text{FPR} \leq 0.005$) without sacrificing accuracy, achieving a false-positive rate of essentially zero and a false-negative rate between 0.5% and 3.8% depending on the threshold and model. We flag a document as a “hit” when the Pangram API determines that its contents has a strictly positive probability of having been generated by an AI tool.

5.2. Sample and method

We draw a stratified random sample of 200 federal civil complaints (either *pro se* or represented) from each calendar year between 2019 and 2026, for a total of 1,600 documents. The sample comes from the CourtListener/RECAP archive, which mirrors the subset of PACER docket filings that contributors have uploaded.¹⁸ Our procedure proceeds as follows:

Randomization. Within each year, we randomly select a day and then a random civil case from that day (`docketNumber:cv`).

Within case. Within each case, there are many documents—these documents can be complaints, motions, or judge memos, but they can also be empty, boilerplate logistical orders, or broken text. To standardize document selection, for each case docket, we choose the *first* document on the docket `entry_number:1`. In federal civil practice, the first docket entry is almost always the initiating complaint. While sampling, for all the days where the docket is found to have a valid RECAP document, we also apply three deterministic filters.

1. First, we require a positive description match to “complaint” or “petition” (standard language on complaint documents) to exclude misfiled or mislabeled entries.

¹⁸We assess that this sample is biased towards represented cases because *pro se* cases are underrepresented in the universe of available case text data. Our procedure targets randomization over the available sample. More details follow through Section 5.3.

2. Second, we filter to non-prisoner NOS codes (for consistency with our sample). This constraint did not bind for any of the documents we sampled.
3. Third, we require that PDFs have at least 500 characters of text (CourtListener applies no quality filtering of their own, and occasionally PDFs are empty beyond a boilerplate PACER header).

If the case has an available RECAP document, and if the document passes these deterministic filters, then we call the complaint ‘valid’ and retain the PDF.

Pangram. Before passing the PDF to Pangram, we strip the PDF of its automatically generated boilerplate PACER header (court name, filing date, case caption)—this boilerplate is not in the complaint as courts read them, and contains no legal content as pertains to the substance of the case. After stripping boilerplate, we pass the complete complaint to the Pangram API for AI-text detection.

Repeat. For each year, across 2019–2026 (four pre-period years, four post-period years), we repeat this process 200 times.

Below, we present the results of this procedure.

5.3. Results

The pre-AI period (2019–2022) yields 1 detection across 800 documents, a rate of 0.1%. The post-AI period (2023–2026) yields 66 detections across 800 documents, a rate of 8.2%. The rate rises monotonically from year to year: 1.0% in 2023, 3.5% in 2024, 10.5% in 2025, and 18.0% in early 2026. Figure 10 plots the share of flagged complaints by year.

The pre-AI baseline is essentially zero, and is conducted to calibrate Pangram’s false positive rate on legal complaints. A single detection out of 800 pre-period complaints means that the detector is not producing false positives at a rate that would confound the post-period signal. In the post-period, we find that the AI detection rate rises consistently, from near zero at the end of 2022 to 18% by early 2026, tracking the trajectory of LLM capability gains and diffusion of LLM adoption rather than any single event. By early 2026, roughly one in five complaint filings contains text that the detector classifies as AI-generated.

That AI use is this visible at the document level, even in a sample that includes attorney-

AI-generated text detection rate, pre- and post-AI

Pooled

Share with any AI detected (%)

25

20

15

10

5

0

Pre-AI

Post-AI

0.1%
(1/800)

8.2%
(66/800)

By year

Share with any AI detected (%)

25

20

15

10

5

0

2019

2020

2021

2022

2023

2024

2025

2026

ChatGPT

Figure 10. Share of federal civil complaints classified as containing AI-generated text, by filing year. Error bars are Wilson 95% confidence intervals. Sample: 200 random complaints per year (2019–2026) drawn from CourtListener/RECAP. Source: Pangram AI-text detector applied to body of complaint texts.

drafted pleadings, is itself notable.¹⁹ Combined with the patterns in Sections 3 and 4, the evidence supports a simple reading: AI is being used to generate federal civil complaints, the use is growing quickly, and it is showing up in filing behavior and in the text of the filings themselves.

An important point of note is that, though we try to maintain a high-quality sample and randomize as completely as possible, there is one important source of bias in the sample. Namely, the universe of RECAP documents available to us is itself not the complete universe of case documents (though it is the most complete sample publicly available). This is because documents are automatically uploaded by users of the RECAP browser extension, disproportionately attorneys, journalists, and legal researchers. The effect is that *pro se* text data is underrepresented. If *pro se* litigants are using AI at a higher rate than attorneys, that suggests that our results are a lower bound for the true prevalence of AI-text in legal cases.²⁰ This is also why it is difficult to complete the same Pangram procedure directly on

¹⁹Of note is that federal judges are prohibited from using AI for drafting opinions under the Judicial Conference’s 2024 interim guidance, and most district courts now require attorneys to certify whether AI was used in any filing. It is difficult to impute to what degree these constraints bind.

²⁰We could only conceive of one way in which *pro se* documents might be over-represented in the post-AI period: if commercially available LLMs being used by *pro se* litigants *explicitly* encourage their users to

pro se cases.

6. Avenues for Future Work

As AI tools continue to advance in capability and diffuse in adoption, we expect that many more hypotheses will become testable. Already, we anticipate that further work along the agenda detailed below is within reach.

First, and most pressingly, we expect that court burden from *pro se* litigants is even higher in state and municipal courts than in federal courts. For most people, if they interact with the judicial system at all, it is likely to be in “lower-stakes” contexts than federal civil litigation. The vast majority (over 90%) of legal cases in the United States are overseen by state and municipal courts which hear family, traffic, and many business matters (National Center for State Courts, 2015). We have already begun expanding our empirical analysis with existing state data and are in search of further data resources to continue.

Second, we also intend to look deeper “inside” each case. Case-level text data is available at varying levels of quality (and coverage) at the federal level. While the results of this paper consider aggregate statistics, textual data also permit us to probe qualitative characteristics of *pro se* cases. Indeed, using advances in natural language processing and LLM-based tools ourselves, we intend to examine whether post-AI complaints are of higher “quality” along various measures: responsiveness to statute, reference to past case law, and the extent of citation hallucination among others. The law progresses a case at a time, and an additional important concern is viewed along these natural language processing abilities in aggregate: are cases homogenizing under shared tool use, and if so, how does that affect courts in the long run?

Third, in addition to examining notions of case quality, we can also use within-case data to better characterize where in the “legal-services production function” AI is having an effect. Consider, for example, the life cycle of a case as a series of production steps toward arriving at a ruling: drafting complaints, filing motions, making arguments, going to (jury) trial, and ultimately writing an opinion. At each of these steps the parties to the case (and the judge) are required to engage in tasks that may each be substituted for or augmented by LLMs.

upload AI-drafted complaints to CourtListener or to install the RECAP browser extension. In anecdotal testing, we were unable to elicit this kind of behavior in a number of widely used AI tools.

Using this approach may not only help illustrate whether LLMs have a factor bias, but also help answer some of the questions raised by the results of this paper. For example, between Sections 4.1 and 4.2 we note that the burden on courts has risen, though this has yet to generate sufficient congestion to have an effect on case resolution times. Using within-docket filing timelines, we intend to test whether the advent of AI tools has accelerated the time between filings to offset the increase in filings.

Fourth, cases are indexed by names of claimants and representation, but our metadata only includes last names, making it difficult to disambiguate between collisions of popular last names. However, we are in the process of cleaning a new source of metadata and matching it to our existing corpus. This will enable us to look more closely at two interesting distributional effects:

1. Our *pro se* results are framed as extensive margin results, but there is also an important intensive margin to consider: whether existing *pro se* filers are now increasing the rate at which they're filing *pro se* complaints. Analysis (based purely on the last names available to us, completed but not presented in this draft) suggests there is an intensive margin effect but that it is small relative to the extensive margin effect. In the coming month, we will be able to state this effect more precisely.
2. Additionally, with the additional docket-level metadata we are enriching, we can extract the identities of the lawyers representing clients. We intend to build a panel of lawyer activity to examine how the advent of AI has changed the distribution and nature of the legal profession. If AI is a general purpose technology that can be adopted by lawyers as well as *pro se* litigants (preliminary evidence for which has been presented in Section 4.2), its effects on lawyer outcomes are *ex ante* ambiguous. In some cases, it could be that the substitution of newly self-represented litigants has in effect replaced these lawyers. Alternatively AI tools could be driving prospective litigants to seek out these lawyers if they are insufficiently confident about those tools to represent themselves. Finally, it may be possible that certain, technically-savvy lawyers are disproportionately augmented by improvements in technology.

Importantly, the latter point has important implications for the industry at large. If AI is substituting for legal professionals for “simple” cases, but labor-augmenting for the most “complex” cases (perhaps because these cases are less rote, or because litigation strategy

becomes more important), then the constituents of the legal industry may benefit or lose differentially (with even more muddled upstream ramifications for prospective lawyers still in law school). Perhaps what has traditionally been a fragmented industry will consolidate, as only the biggest firms can offer services that AI tools cannot compete with. Alternatively, perhaps AI tools empower independent practitioners to take on ever more complex cases and demanding clients (for evidence of labor market effects during a previous episode of automation, see e.g. Feigenbaum and Gross (2024)). We do not take a stance on which effects will dominate or how the legal industry will respond to increasingly capable generative AI tools. However, we anticipate that the labor market for lawyers and other legal professionals will change in the months and years to come and we expect to document it closely.

7. Conclusion

The federal civil courts are absorbing a large, rapid, technology-driven increase in *pro se* filings. That the *pro se* activity is happening in federal court, rather than in the more accessible and lower-stakes state or municipal courts, is itself striking. Federal civil cases require satisfying federal jurisdictional rules, surviving federal pleading standards, and navigating a procedural apparatus that is widely acknowledged to be among the most demanding in the American legal system. If generative AI is enabling self-representation in this demanding context, the effect in state and local courts is likely larger still.

Within the federal system, the supply side of the adjustment margin is tightly constrained. Article III judges require years of training and Senate confirmation, and the authorized number of district judgeships has grown only modestly since the last comprehensive expansion in 1990 (Menell and Vacca, 2021). Current Judicial Conference guidance prohibits federal judges from using AI to draft opinions, and most district courts require attorneys to certify whether AI was used in any filing. The supply of adjudicative capacity is therefore fixed on nearly every margin: the number of judges is fixed by Congress, the training time of new judges is fixed by the profession, and the ability of existing judges to use AI to process more cases is prohibited by policy. This paper has documented a rise in demand for adjudicative services. It is hard to see how the mismatch resolves without some form of adjustment.

Beyond the direct concerns of court burden, two possible trajectories are worthy of note. The first is an arms race between private litigants. If represented plaintiffs are already

responding to rising *pro se* filing activity by generating more written submissions of their own, the post-AI equilibrium may feature cases where docket-activity counts drift upward on both sides, absorbing an increasing share of judicial attention per case without resolving a correspondingly greater share of disputes. The per-case entries increases we document by 2025Q2—38% for *pro se* and 23% for represented—are already meaningful, and small relative to what this process could produce if it continues for years to come.

The second is asymmetric litigation against institutional defendants, particularly the federal government and large administrative agencies that cannot easily scale their litigation capacity in response. A federal agency sued by a plaintiff with AI-drafted filings faces an adversary whose marginal filing cost has fallen, while the agency’s own response cost has not. The civil rights, immigration, and benefits-appeal categories in which we find the largest *pro se* increases are disproportionately cases against federal defendants. If AI scales the rate at which such cases are filed faster than it scales the rate at which agencies can respond, the government accumulates a queue of cases that it cannot process, with consequences for administration as well as adjudication.

The structural question is whether the supply of adjudicative services can grow at all. One margin for this is to relax the rules that prohibit federal judges from using AI assistance, enabling a productivity response on the adjudication side analogous to the one AI has already produced on the filing side. Another is to more aggressively route simple cases—the cases we show drive the bulk of the post-AI rise—to magistrate judges, lower courts, or specialized triage procedures. If lower courts are already burdened, this may even require more drastic action, as in the construction of a new lower court that triages the simplest cases mostly with AI. These avenues are not mutually exclusive, and neither is costless. Both may be preferable to the alternative, which is an accumulation of cases in a system whose throughput is fixed.

Generative AI has lowered the cost of producing passable legal documents. Americans are noticing, and they are walking into federal civil court in much larger numbers to act on it. The courts are not yet visibly buckling under the load—case durations are stable and dispositions look familiar—but the volume of activity inside cases has risen sharply and shows no signs of stopping. Whether the federal civil court system and the broader legal community absorbs this technology as productivity gain or as congestion is being determined in real time.

References

- [1] Daron Acemoglu. “The simple macroeconomics of AI”. In: *Economic Policy* 40.121 (2025), pp. 13–58.
- [2] Administrative Office of the U.S. Courts. *Federal Judicial Caseload Statistics: Judicial Caseload Indicators*. Tech. rep. Administrative Office of the U.S. Courts, 2024. URL: <https://www.uscourts.gov/data-news/reports/statistical-reports/federal-judicial-caseload-statistics/judicial-caseload-indicators-federal-judicial-caseload-statistics-2024>.
- [3] Philippe Aghion and Simon Bunel. “AI and growth: Where do we stand?” In: *policy note* (2024).
- [4] Alberto Alesina and Eliana La Ferrara. “A Test of Racial Bias in Capital Sentencing”. In: *American Economic Review* 104.11 (2014), pp. 3397–3433.
- [5] Lukas Althoff and Hugo Reichardt. “Task-Specific Technical Change and Comparative Advantage”. In: *CESifo Working Paper* (2026).
- [6] David Arnold, Will Dobbie, and Crystal S. Yang. “Racial Bias in Bail Decisions”. In: *Quarterly Journal of Economics* 133.4 (2018), pp. 1885–1932.
- [7] Elliott Ash, Daniel L. Chen, and Suresh Naidu. “Ideas Have Consequences: The Impact of Law and Economics on American Justice”. In: *Quarterly Journal of Economics* (2026). Forthcoming.
- [8] Gary S Becker. “A Theory of the Allocation of Time”. In: *The economic journal* 75.299 (1965), pp. 493–517.
- [9] Erik Brynjolfsson, Bharat Chandar, and Ruyu Chen. *Canaries in the Coal Mine? Six Facts about the Recent Employment Effects of Artificial Intelligence*. Tech. rep. Stanford Digital Economy Lab, 2025. URL: <https://digitaleconomy.stanford.edu/publications/canaries-in-the-coal-mine/>.
- [10] Erik Brynjolfsson, Danielle Li, and Lindsey R. Raymond. “Generative AI at Work”. In: *NBER Working Paper* 31161 (2023).
- [11] Erik Brynjolfsson, Daniel Rock, and Chad Syverson. “The Productivity J-Curve: How Intangibles Complement General Purpose Technologies”. In: *American Economic Journal: Macroeconomics* 13.1 (2021), pp. 333–372.

- [12] Alma Cohen and Crystal S. Yang. “Judicial Politics and Sentencing Decisions”. In: *American Economic Journal: Economic Policy* 11.1 (2019), pp. 160–191.
- [13] Robert Collinson, John Eric Humphries, Nicholas Mader, Davin Reed, Daniel Tannenbaum, and Winnie van Dijk. “Eviction and Poverty in American Cities”. In: *Quarterly Journal of Economics* 139.1 (2024), pp. 57–120.
- [14] Fabrizio Dell’Acqua, Edward McFowland, Ethan R. Mollick, Hila Lifshitz-Assaf, Katherine Kellogg, Saran Rajendran, Lisa Kraye, François Candelon, and Karim R. Lakhani. “Navigating the Jagged Technological Frontier: Field Experimental Evidence of the Effects of AI on Knowledge Worker Productivity and Quality”. In: *Harvard Business School Working Paper* 24-013 (2023).
- [15] Theodore Eisenberg and Kevin M. Clermont. “Plaintiphobia in the Supreme Court”. In: *Cornell Law Review* 100 (2014), pp. 193–230.
- [16] Tyna Eloundou, Sam Manning, Pamela Mishkin, and Daniel Rock. “GPTs are GPTs: Labor Market Impact Potential of LLMs”. In: *Science* 384.6702 (2024), pp. 1306–1308. DOI: [10.1126/science.adj0998](https://doi.org/10.1126/science.adj0998).
- [17] Federal Judicial Center. *Integrated Database: Civil Cases Filed in U.S. District Courts*. Tech. rep. Available at <https://www.fjc.gov/research/idb>. Federal Judicial Center, 2025.
- [18] James Feigenbaum and Daniel P Gross. “Answering the call of automation: How the labor market adjusted to mechanizing telephone operation”. In: *The Quarterly Journal of Economics* 139.3 (2024), pp. 1879–1939.
- [19] Mark D. Gough and Emily Taylor Poppe. “(Un)Changing Rates of Pro Se Litigation in Federal Court”. In: *Law & Social Inquiry* 45.3 (2020), pp. 567–596.
- [20] D. James Greiner, Cassandra Wolos Pattanayak, and Jonathan Hennessy. “The Limits of Unbundled Legal Assistance: A Randomized Study in a Massachusetts District Court and Prospects for the Future”. In: *Harvard Law Review* 126.4 (2013), pp. 901–989.
- [21] Neel Guha, Julian Nyarko, Daniel E. Ho, Christopher Ré, Adam Chilton, Aditya Narayana, Alex Chohlas-Wood, Austin Peters, Brandon Waldon, Daniel N. Rockmore, Diego Zambrano, Dmitry Talisman, Enam Hoque, Faiz Surani, Frank Fagan, Galit Sarfaty, Gregory M. Dickinson, Haggai Porat, Jason Hegland, Jessica Wu, Joe Nudell, Joel Niklaus, John Nay, Jonathan H. Choi, Kevin Tobia, Margaret Hagan, Megan Ma, Michael Livermore, Nikon Rasumov-Rahe, Nils Holzenberger, Noam Kolt, Peter

- Henderson, Sean Rehaag, Sharad Goel, Shang Gao, Spencer Williams, Sunny Gandhi, Tom Zur, Varun Iyer, and Zehua Li. “LegalBench: A Collaboratively Built Benchmark for Measuring Legal Reasoning in Large Language Models”. In: *Advances in Neural Information Processing Systems 36 (NeurIPS 2023) Datasets and Benchmarks Track*. 2023.
- [22] Gillian K. Hadfield. *Rules for a Flat World: Why Humans Invented Law and How to Reinvent It for a Complex Global Economy*. Oxford University Press, 2016.
- [23] Alex Imas and Brian Jabarian. *Detecting AI-Generated Text*. Working Paper 34223. National Bureau of Economic Research, 2025. URL: <https://www.nber.org/papers/w34223>.
- [24] Andrew Johnston and Christos Makridis. “AI, Output, and Employment”. In: *CESifo Working Paper* (2026).
- [25] Daniel Martin Katz, Michael James Bommarito, Shang Gao, and Pablo Arredondo. “GPT-4 Passes the Bar Exam”. In: *Philosophical Transactions of the Royal Society A* 382.2270 (2024). Working paper available at SSRN 4389233 (2023).
- [26] Jon Kleinberg, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan. “Human Decisions and Machine Predictions”. In: *Quarterly Journal of Economics* 133.1 (2018), pp. 237–293.
- [27] Alexandra D. Lahav and Peter Siegelman. “The Curious Incident of the Falling Win Rate”. In: *UC Davis Law Review* 52.3 (2019), pp. 1371–1430.
- [28] Legal Services Corporation. *The Justice Gap: The Unmet Civil Legal Needs of Low-Income Americans*. Legal Services Corporation, 2022.
- [29] Mitchell N. Levy. “Empirical Patterns of *Pro Se* Litigation in Federal District Courts”. In: *University of Chicago Law Review* 85.7 (2018), pp. 1819–1870.
- [30] Peter S. Menell and Ryan G. Vacca. “Revisiting and Confronting the Federal Judiciary Capacity “Crisis”: Charting a Path for Federal Judiciary Reform”. In: *California Law Review* 109.4 (2021), pp. 1211–1316.
- [31] National Center for State Courts. *The Landscape of Civil Litigation in State Courts*. Tech. rep. National Center for State Courts, 2015. URL: https://www.ncsc.org/__data/assets/pdf_file/0020/13376/civiljusticereport-2015.pdf.
- [32] Shakked Noy and Whitney Zhang. “Experimental Evidence on the Productivity Effects of Generative Artificial Intelligence”. In: *Science* 381.6654 (2023), pp. 187–192.

- [33] Nicholas G. Otis, Rowan Clarke, Solène Delecourt, David Holtz, and Rembrand Koning. “The Uneven Impact of Generative AI on Entrepreneurial Performance”. In: *Harvard Business School Working Paper 24-042* (2024). SSRN:4671369.
- [34] Deborah L. Rhode. “Access to Justice”. In: *Fordham Law Review* 69.5 (2004), pp. 1785–1819.
- [35] Rebecca L. Sandefur. “Access to What?” In: *Daedalus* 148.1 (2019), pp. 49–55.
- [36] Suproteem K Sarkar. “Ai agents, productivity, and higher-order thinking: Early evidence from software development”. In: *Available at SSRN 5713646* (2025).
- [37] Jessica K. Steinberg. “Demand Side Reform in the Poor People’s Court”. In: *Connecticut Law Review* 47.3 (2023), pp. 741–810.
- [38] Lee C Tucker. “You’re (not) hired: Artificial intelligence and early career hiring in the Quarterly Workforce Indicators”. In: *Working Paper* (2026).
- [39] Parker Whitfill, Cheryl Wu, Joel Becker, and Nate Rush. *Many SWE-bench-Passing PRs Would Not Be Merged into Main*. <https://metr.org/notes/2026-03-10-many-swe-bench-passing-prs-would-not-be-merged-into-main/>. Mar. 2026.
- [40] Emma Wiles, Zanele Munyikwa, and John J. Horton. “Algorithmic Writing Assistance on Jobseekers’ Resumes Increases Hires”. In: *Management Science* 71.12 (2025). NBER Working Paper 30886.
- [41] Jefri Wood. *Pro Se Case Management for Nonprisoner Civil Litigation*. Tech. rep. Federal Judicial Center, 2016. URL: <https://www.fjc.gov/content/315899/pro-se-case-management-nonprisoner-civil-litigation>.

A. The Vermont outlier: immigration mandamus filings

The District of Vermont is the single state-level outlier in Figure 5. Pro se filings in Vermont rose from roughly 45 cases per year before 2022 to over 1,100 in FY2024. Nearly all of the increase comes from a single NOS category—465, “Other Immigration Actions”—and a single type of claim: writs of mandamus filed against the Director of United States Citizenship and Immigration Services (USCIS). We describe the phenomenon here and show that excluding Vermont from the sample leaves the paper’s results qualitatively unchanged.

Background

The USCIS is the federal agency responsible for processing green-card applications, naturalization petitions, and other immigration benefits. In recent years, processing times for several categories of applications have lengthened substantially, with some applicants waiting two or more years for routine adjudications. A writ of mandamus is a court order compelling a government official to perform a duty owed to the petitioner. When USCIS delays an application beyond what the applicant considers reasonable, the applicant can file a mandamus action in federal court naming the USCIS Director as defendant, asking the court to order USCIS to adjudicate the application.

The District of Vermont became a preferred venue for these filings. It is a small court with a fast docket, and the named defendant—Ur Mendoza Jaddou, Director of USCIS—could be sued there under permissive venue rules. Beginning in 2023, online communities of immigration applicants began sharing detailed instructions for filing mandamus actions *pro se* in the District of Vermont, explicitly leveraging AI tools to draft the filings.

A representative guide

Figure A1 reproduces excerpts from a widely circulated Reddit post offering step-by-step instructions for filing a writ of mandamus without a lawyer. The guide’s first step reads: “I used AI (Microsoft Copilot, to be exact). It’s free and comes with Windows computers. I asked it to write me a writ of mandamus. I knew just AI wasn’t enough, and I know nothing about law, so I did step 2.” Step 2 describes hiring a Fiverr lawyer for \$150 to review and polish the AI-generated draft. Step 3 lists the named defendants: USCIS, the

Attorney General, the USCIS Director, the Secretary of Homeland Security, and the local U.S. Attorney.

Step-by-Step Guide on How to File a Writ of Mandamus by Yourself:

Step 1: I used AI (Microsoft Copilot, to be exact). It's free and comes with Windows computers. I asked it to write me a writ of mandamus. I knew just AI wasn't enough, and I know nothing about law, so I did step 2.

Step 2: I went on Fiverr and messaged several immigration lawyers. I found one who was very knowledgeable. I told her I made a writ of mandamus and needed her to edit it and make it worthy enough to take to court. She charged me \$150. (If you want to double-check with another Fiverr lawyer after your own edits, it should not cost you more than \$150—some people tried to charge me thousands.) She did an amazing job. I did a few more edits after her, and it was perfect. I will include it at the end redacted. I'll even mark places where you need to put your personal info and what to do. Feel free to make edits if you like.

Step 3: On the top of the writ of mandamus, you will see the defendants. The following people are included as defendants:

- United States Citizenship and Immigration Services (USCIS)
- Merrick Garland, Attorney General of the United States
- Ur Mendoza Jaddou, Director of USCIS
- Alejandro Mayorkas, Secretary of the Department of Homeland Security
- Markenzy Lapointe, United States Attorney for the Southern District of Florida

Figure A1. Excerpts from a Reddit post (r/USCIS, December 2024) providing step-by-step instructions for filing a *pro se* writ of mandamus using AI and a \$150 Fiverr lawyer review.

The Reddit post suggests that applicants with no legal training use an LLM to generate an initial draft, pay a small amount for a lawyer to review it, and file the result in a strategically chosen venue. These are legitimate individual cases—each petitioner has a specific pending application and a specific grievance about processing time—but the volume is enabled by the dramatically lower cost of producing the necessary legal documents.

Removing Vermont from the sample

Vermont's contribution to the national *pro se* count is small in absolute terms. Table 2 reports Vermont *pro se* case counts and their effect on the national IDB *pro se* share.

At its peak impact (FY2024), excluding Vermont reduces the national *pro se* share by 0.48 percentage points—from 17.09% to 16.61%. The national rise from the pre-AI baseline of roughly 14% to the FY2025 level of 19.6% is essentially unchanged. Figure A2 plots the national headline shares and counts with and without Vermont; the two series are nearly

Table 2. Vermont *pro se* cases and impact on national IDB *pro se* share.

FY	Vermont PS		National IDB PS share		
	Cases	× pre-AI	With VT	Without VT	Difference
2019	43	1.0×	14.09%	14.09%	<0.01 pp
2020	47	1.0×	12.60%	12.59%	0.01 pp
2021	45	1.0×	12.48%	12.47%	0.01 pp
2022	51	1.1×	14.09%	14.08%	0.01 pp
2023	307	6.8×	14.94%	14.82%	0.13 pp
2024	1,121	24.9×	17.09%	16.61%	0.48 pp
2025	1,018	22.6×	19.92%	19.55%	0.37 pp

indistinguishable.

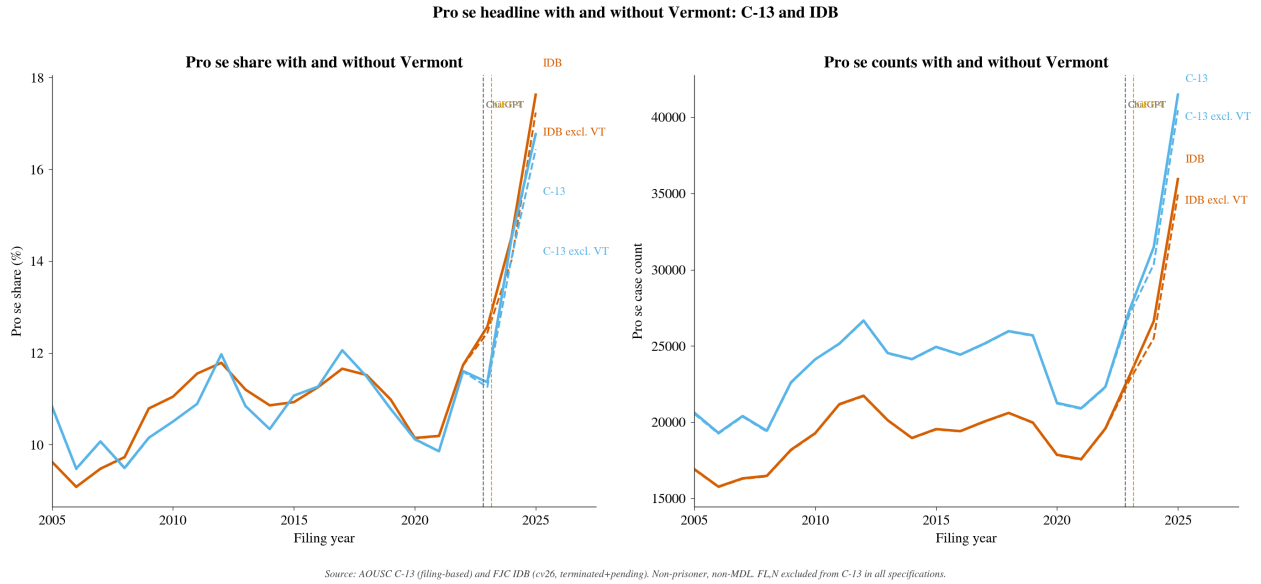


Figure A2. National *pro se* filing counts (left) and share (right) with and without the District of Vermont. The two series are nearly identical. Source: FJC IDB, non-prisoner non-MDL.

B. Additional Figures and Exhibits

B.1. Simple/complex partition: NOS code listing

Table 3 lists all 86 non-prisoner NOS codes included in the analysis, with pre-AI average annual case volume, pre-AI *pro se* share, and partition assignment. The partition is determined by the case-weighted median of the pre-AI *pro se* share across all 86 codes; codes above the median are assigned to “Simple” and codes below to “Complex.” All pre-AI statistics are

computed over FY2015–FY2022. Within each partition, codes are sorted by pre-AI *pro se* share (descending).

Table 3. NOS codes, pre-AI case volume, pre-AI *pro se* share, and partition assignment. Pre-AI statistics computed over FY2015–FY2022. Partition split at the case-weighted median pre-AI *pro se* share.

NOS	Description	Pre-AI cases	Pre-AI PS%	Partition
<i>Simple (above case-weighted median pre-AI pro se share)</i>				
230	Rent Lease & Ejectment	6,534	87.1%	Simple
463	Habeas Corpus – Alien Detainee	9,591	54.7%	Simple
400	State Reapportionment	41	51.2%	Simple
440	Other Civil Rights	114,998	50.7%	Simple
153	Recovery of Overpayment of Vet. Benefits	49	46.9%	Simple
320	Assault, Libel & Slander	4,548	44.8%	Simple
371	Truth in Lending	1,776	43.9%	Simple
460	Deportation	105	41.0%	Simple
290	All Other Real Property	12,327	36.1%	Simple
430	Banks and Banking	1,373	35.6%	Simple
422	Appeal 28 USC 158	14,097	35.3%	Simple
950	Constitutionality of State Statutes	2,969	33.8%	Simple
441	Voting	1,094	30.7%	Simple
220	Foreclosure	21,568	30.5%	Simple
470	RICO	6,293	30.5%	Simple
871	IRS – Third Party 26 USC 7609	286	29.0%	Simple
870	Taxes (US Plaintiff or Defendant)	5,551	25.4%	Simple
443	Housing/Accommodations	6,603	25.1%	Simple
370	Other Fraud	16,086	21.7%	Simple
423	Withdrawal 28 USC 157	2,155	21.2%	Simple
140	Negotiable Instrument	3,121	19.0%	Simple
862	Black Lung (923)	37	18.9%	Simple
362	Personal Injury – Medical Malprac- tice	8,987	18.6%	Simple
485	Telephone Consumer Protection Act	3,410	18.6%	Simple
890	Other Statutory Actions	57,797	18.0%	Simple
442	Employment	92,089	18.0%	Simple
375	False Claims Act	3,335	17.0%	Simple

Continued on next page

Table 3 continued

NOS	Description	Pre-AI cases	Pre-AI PS%	Partition
893	Environmental Matters	6,395	16.2%	Simple
861	HIA (1395ff)	669	15.4%	Simple
445	Amer. w/Disabilities – Employment	18,824	14.9%	Simple
895	Freedom of Information Act	5,638	14.5%	Simple
152	Recovery of Defaulted Student Loans	3,522	14.3%	Simple
380	Other Personal Property Damage	9,015	14.3%	Simple
376	Qui Tam (31 USC 3729a)	3,599	14.0%	Simple
899	Admin. Procedure Act/Review of Agency Decision	6,956	13.9%	Simple
625	Drug Related Seizure of Property	4,340	13.5%	Simple
450	Commerce	1,316	13.2%	Simple
448	Education	5,513	13.1%	Simple
150	Recovery of Overpayment	3,500	12.2%	Simple
240	Torts to Land	3,004	12.1%	Simple
690	Other	4,230	12.1%	Simple
740	Railway Labor Act	408	10.8%	Simple
490	Cable/Sat TV	4,237	10.6%	Simple
190	Other Contract	83,511	10.3%	Simple
360	Other Personal Injury	71,540	10.2%	Simple
410	Antitrust	3,283	9.6%	Simple
880	Customs	1,482	9.3%	Simple
480	Consumer Credit	79,869	9.2%	Simple
<i>Complex (below case-weighted median pre-AI pro se share)</i>				
462	Naturalization Application	2,658	9.0%	Complex
891	Agricultural Acts	1,341	8.9%	Complex
210	Land Condemnation	2,532	8.4%	Complex
720	Labor/Management Relations	4,274	8.4%	Complex
850	Securities/Commodities/Exchange	12,745	8.1%	Complex
151	Medicare Act	1,323	8.1%	Complex
355	Motor Vehicle Product Liability	2,194	7.8%	Complex
345	Marine Product Liability	103	7.8%	Complex
196	Franchise	2,080	7.6%	Complex
864	SSID Title XVI	63,359	7.5%	Complex
896	Arbitration	6,466	6.9%	Complex

Continued on next page

Table 3 continued

NOS	Description	Pre-AI cases	Pre-AI PS%	Partition
790	Other Labor Litigation	10,861	6.7%	Complex
160	Stockholders Suits	2,146	6.7%	Complex
840	Trademark	24,849	6.6%	Complex
368	Asbestos Personal Injury Product Liability	1,355	6.6%	Complex
330	Federal Employers Liability	2,849	6.5%	Complex
465	Other Immigration Actions	24,372	6.1%	Complex
863	DIWC/DIWW (405g)	72,822	6.0%	Complex
710	Fair Labor Standards Act	56,331	5.6%	Complex
310	Airplane	1,512	5.6%	Complex
120	Marine	5,567	5.5%	Complex
820	Copyrights	36,415	5.4%	Complex
446	Amer. w/Disabilities – Other	76,874	5.3%	Complex
245	Tort Product Liability	2,250	5.2%	Complex
367	Health Care/Pharmaceutical PI Product Liability	14,780	5.2%	Complex
110	Insurance	92,760	4.8%	Complex
340	Marine (General)	9,176	4.6%	Complex
791	Employee Retirement Income Security Act	46,621	4.2%	Complex
365	Personal Injury – Product Liability	28,297	4.1%	Complex
195	Contract Product Liability	2,947	4.1%	Complex
751	Family and Medical Leave Act	9,548	3.8%	Complex
865	RSI (405g)	6,419	3.2%	Complex
350	Motor Vehicle	38,802	3.2%	Complex
385	Property Damage Product Liability	5,319	3.1%	Complex
315	Airplane Product Liability	509	2.9%	Complex
830	Patent	31,353	2.6%	Complex
130	Miller Act	2,021	2.5%	Complex
835	Patent – Abbreviated New Drug Application	1,585	1.6%	Complex

B.2. Characterizing the simple/complex partition

Table 1 uses only the pre-AI FY2005–FY2022 non-prisoner, non-MDL sample used to construct the paper’s simple/complex partition. We do not infer these buckets from the simple/complex split itself. “Federal question” is $JURIS=3$; “U.S.-party” is $JURIS \in \{1, 2\}$. For the NOS-based characterization, we fix three ex ante buckets shown in Table 4. Review / petition collects NOS codes whose initiating paper asks the district court to review, compel, or correct a prior agency, custody, or bankruptcy decision. Standardized complaint collects NOS codes that map cleanly to national AO civil *pro se* complaint forms.²¹ Specialist complaint is defined narrowly as NOS codes with claim-specific pleading, filing, or bar-admission regimes beyond ordinary notice pleading.²²

Two due-diligence points matter for interpretation. First, the standardized bucket is intentionally conservative. Many NOS codes with high *pro se* rates — consumer credit (480), foreclosure (220), and general contract disputes — lack the claim-specific pleading requirements of the specialist case types, and their corresponding AOUSC forms cover only a subset of the NOS code’s full case population. These codes are nonetheless included in the standardized bucket when a matching AO form exists; the remaining NOS codes with no matching form stay in “other original complaint.” Second, *JURY* does not sharpen the interpretation. If anything, simple-side cases demand juries more often than complex-side cases (52.0% versus 43.2%), largely because civil-rights and employment complaints frequently request jury trials. We therefore use *JURY* as a diagnostic check, not as a defining row in the main-text table.

²¹See the official AO Civil Pro Se Forms page, which includes national complaint forms for employment discrimination, wage-and-hour claims, civil-rights complaints, and Social Security review: uscourts.gov/forms-rules/forms/civil-pro-se-forms. Social Security review (NOS 861–865) has a dedicated AO Form 85 but is classified here as review / petition because the case is fundamentally record review of an agency determination, not an original complaint on the merits.

²²Patent (NOS 830, 835): practitioners must be registered with the USPTO under 37 C.F.R. §11.7. Securities (NOS 850): PSLRA imposes heightened pleading and automatic discovery stay for private actions, 15 U.S.C. §78u-4; *Tellabs, Inc. v. Makor Issues & Rights, Ltd.*, 551 U.S. 308 (2007). Copyright (NOS 820): registration is a prerequisite to suit, *Fourth Estate Public Benefit Corp. v. Wall-Street.com, LLC*, 586 U.S. 296 (2019). RICO (NOS 470): plaintiffs must plead an enterprise plus a pattern of racketeering activity; fraud predicates require particularity under Fed. R. Civ. P. 9(b). FCA / *Qui Tam* (NOS 375, 376): the complaint is filed under seal and served only on the United States, 31 U.S.C. §3730(b)(2), and the government must decide whether to intervene before the action proceeds.

Table 4. NOS-code definitions for the characterization buckets.

Bucket	NOS codes	Institutional rationale
Standardized complaint	440, 441, 443, 444, 448 (civil rights); 442, 445 (employment discrimination); 190 (contract); 360 (other PI); 380 (personal property damage); 220, 230, 240, 290 (real property)	NOS codes with a corresponding national AO civil <i>pro se</i> complaint form. Filing-stage work is well-specified and template-amenable.
Specialist complaint	830, 835 (patent); 850 (securities); 820 (copyright); 470 (RICO); 375, 376 (FCA / <i>qui tam</i>)	Claim-specific pleading, registration, or bar-admission requirements beyond ordinary notice pleading (37 C.F.R. §11.7; 15 U.S.C. §78u-4; <i>Fourth Estate</i> ; 31 U.S.C. §3730(b)(2)).
Review / petition	861–865 (Social Security); 899, 895 (APA review, FOIA); 462, 463, 465 (immigration); 422, 423 (bankruptcy appeal / withdrawal); 870, 871 (tax); 791 (ERISA record review)	The initiating paper asks the district court to review, compel, or correct a prior agency, custody, or bankruptcy decision. The litigation is record-bound rather than original-complaint-driven.
Other original complaint	All remaining NOS codes	Does not map to an AO form, a specialist regime, or a review vehicle.

B.3. Simple/complex partition: ablations across percentile cutpoints

Figure 4 in the main text partitions NOS codes at the case-weighted median (50th percentile) of pre-AI *pro se* share. The following figures replicate the same plot at the 80th, 70th, 60th, 40th, and 30th percentile cutpoints. The separation between the high-*pro se* and low-*pro se* groups is robust across all cutpoints, confirming that the main finding does not depend on the particular choice of threshold.

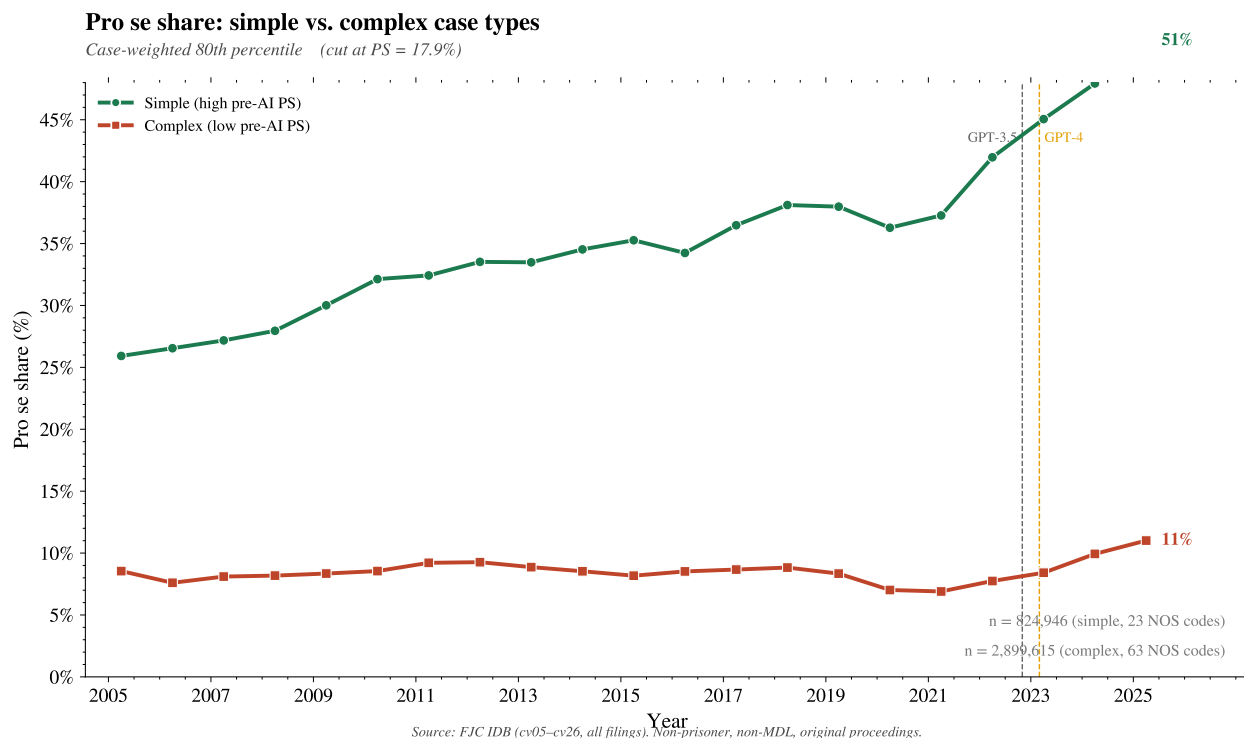


Figure B3. Simple/complex split at the 80th percentile of pre-AI *pro se* share (case-weighted). NOS codes above the 80th percentile are “Simple”; all others are “Complex.” Source: FJC IDB, non-prisoner, non-MDL cases.

B.4. 2025 *pro se* rates compared to pre-AI peak

Constructed using the same method at the district level, there are five districts that see contractions in *pro se* filings within districts that are depicted as having experience an increase in the Figure 5: the Southern District of Illinois, the Western District of Missouri, the Western District of New York, the Eastern District of Oklahoma, and the Southern District of South Dakota. The full listing of district-level *pro se* changes in filings from their pre-AI peaks to FY2025 is presented in Table 5.

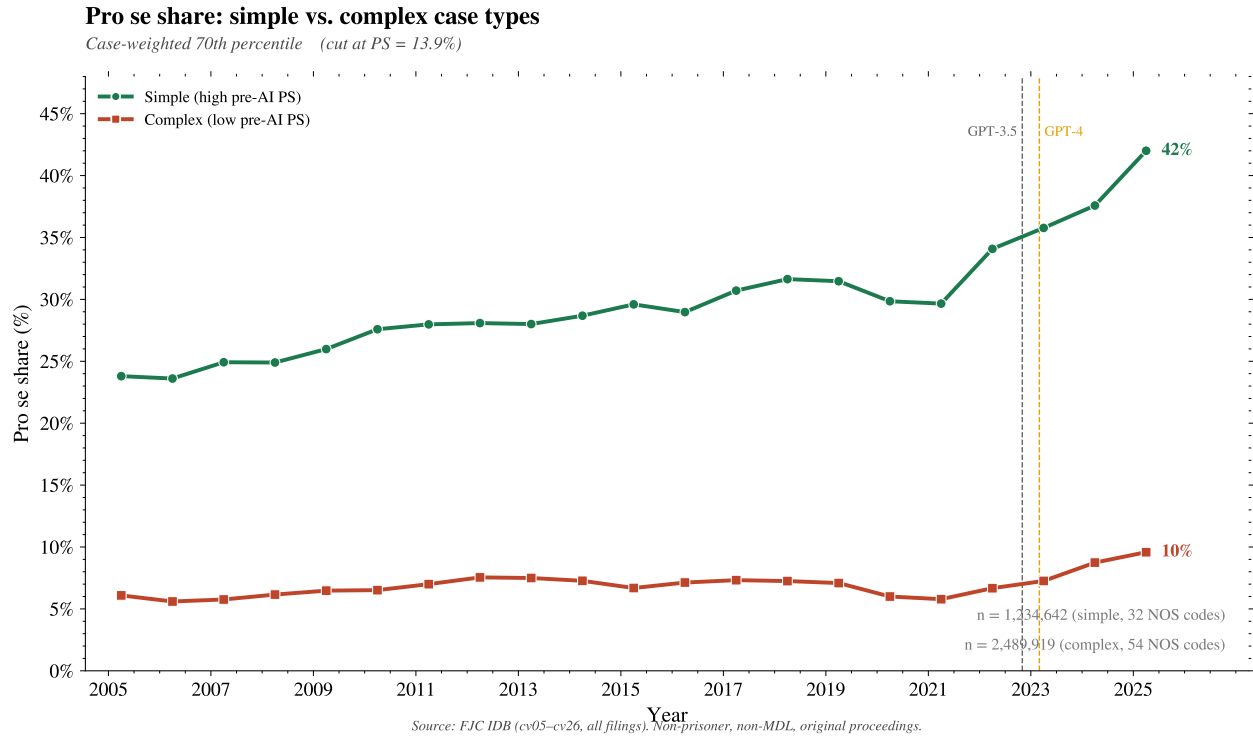


Figure B4. Simple/complex split at the 70th percentile of pre-AI *pro se* share (case-weighted).
Source: FJC IDB, non-prisoner, non-MDL cases.

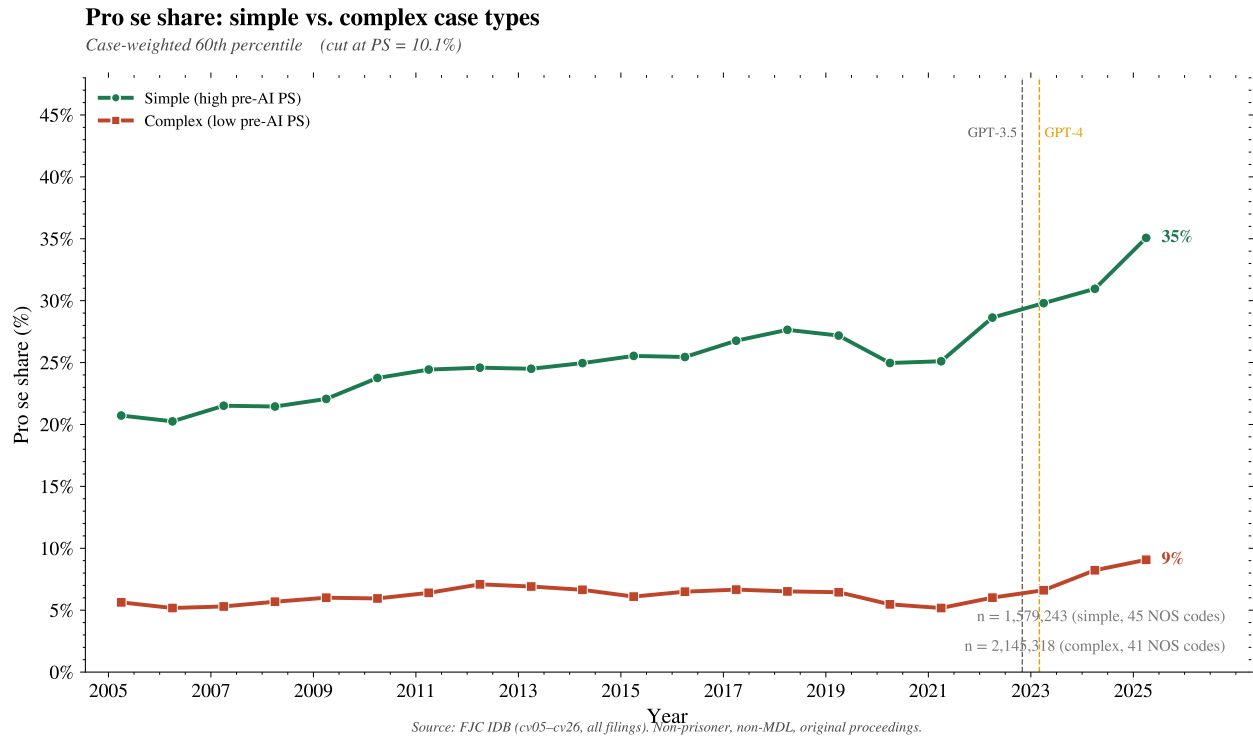


Figure B5. Simple/complex split at the 60th percentile of pre-AI *pro se* share (case-weighted).
Source: FJC IDB, non-prisoner, non-MDL cases.

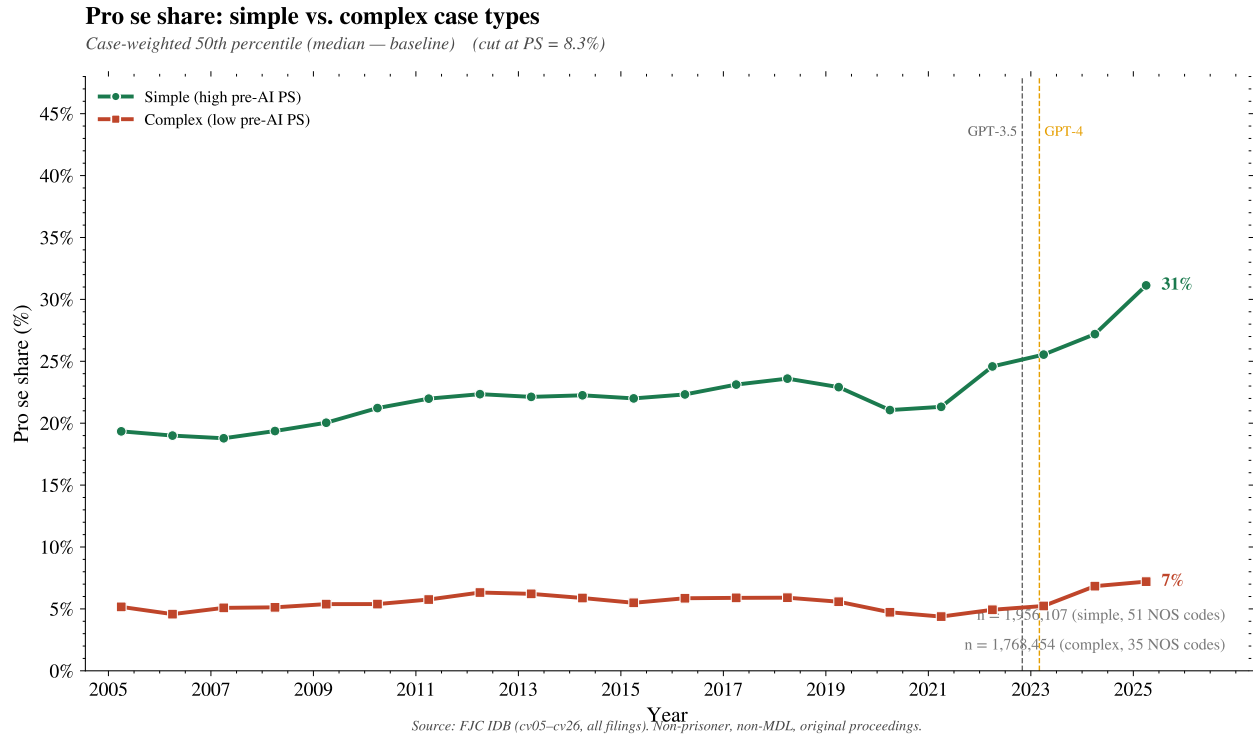


Figure B6. Simple/complex split at the 50th percentile of pre-AI *pro se* share (case-weighted), matching the main-text Figure 4. Source: FJC IDB, non-prisoner, non-MDL cases.

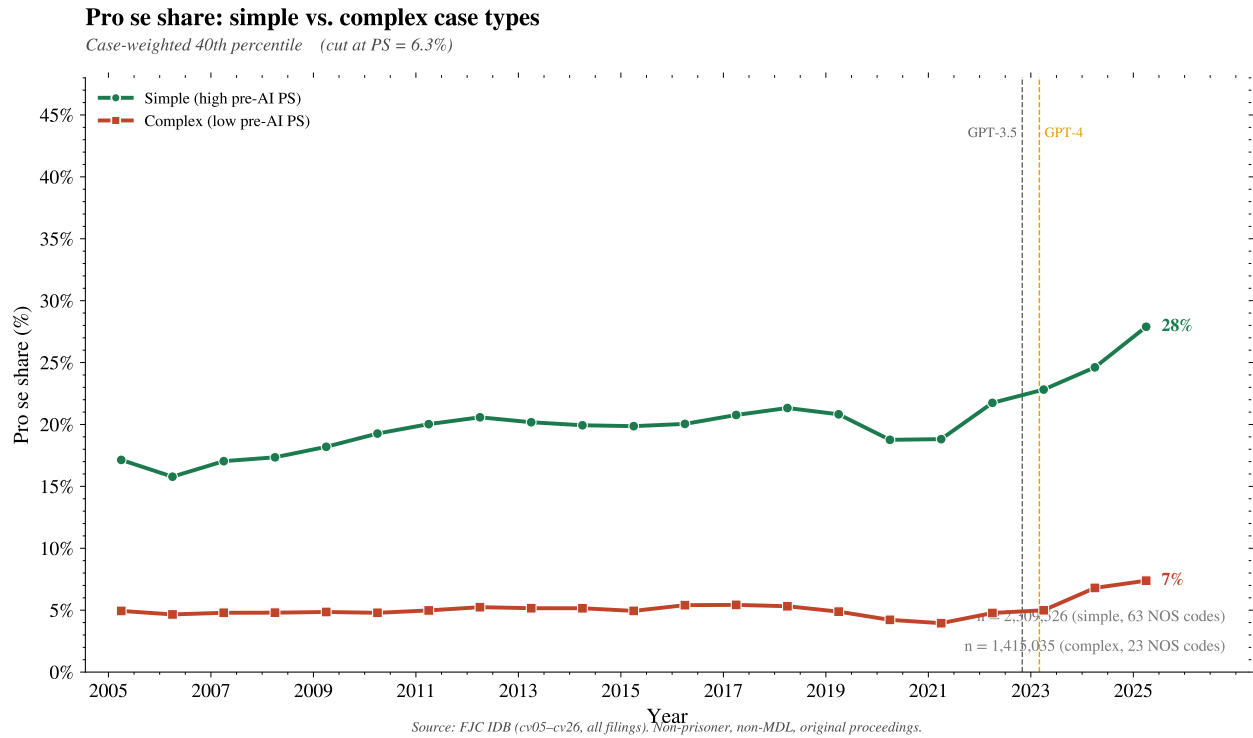


Figure B7. Simple/complex split at the 40th percentile of pre-AI *pro se* share (case-weighted). Source: FJC IDB, non-prisoner, non-MDL cases.

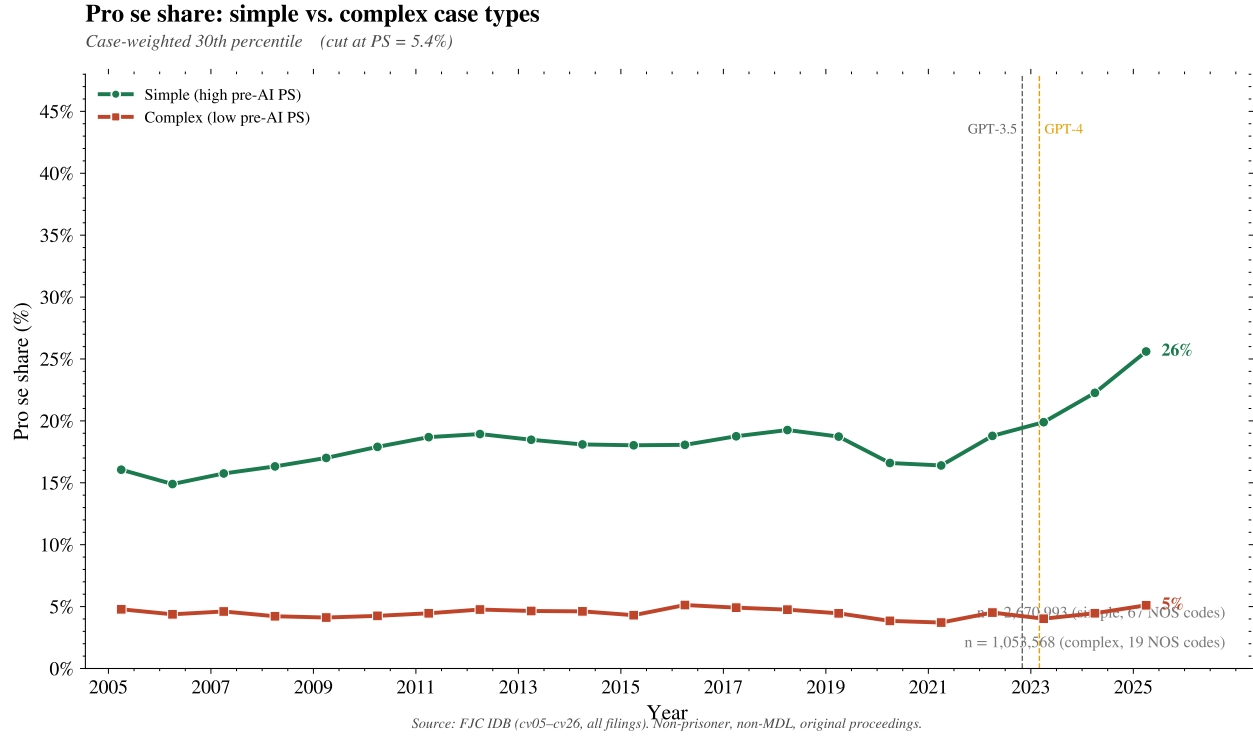


Figure B8. Simple/complex split at the 30th percentile of pre-AI *pro se* share (case-weighted). Source: FJC IDB, non-prisoner, non-MDL cases.

Table 5. District-level FY2025 *pro se* filings relative to each district court’s peak pre-AI year in FY2008–FY2022. Bold rows indicate that the district court saw a *pro se* filing contraction from the pre-AI peak.

District court	Pre-AI peak		% change from pre-AI peak	District court	Pre-AI peak		% change from pre-AI peak
	Pro se count	Year			Pro se count	Year	
AK	519	2010	16.8%	MS - Southern	225	2020	-32.0%
AL - Middle	107	2009	121%	MT	72	2013	-4.17%
AL - Northern	375	2012	-13.3%	NC - Eastern	319	2013	3.45%
AL - Southern	86	2012	-20.9%	NC - Middle	126	2014	38.1%
AR - Eastern	142	2016	15.5%	NC - Western	169	2018	68.6%
AR - Western	80	2012	1.25%	ND	36	2018	86.1%
AZ	183	2022	-30.6%	NE	107	2021	171%
CA - Central	1656	2011	23.7%	NH	87	2019	37.9%
CA - Eastern	512	2012	26.2%	NJ	804	2015	22.6%
CA - Northern	685	2012	22.8%	NM	174	2017	0.00%
CA - Southern	373	2010	5.36%	NV	550	2021	-1.27%
CO	417	2020	48.2%	NY - Eastern	770	2017	9.09%
CT	310	2018	19.7%	NY - Northern	256	2022	12.1%
DC	636	2018	112%	NY - Southern	1150	2015	33.0%
DE	163	2013	48.5%	NY - Western	264	2013	-14.0%

Continued on next page

Table 5 continued from previous page

District court	Pre-AI peak		% change from pre-AI peak	District court	Pre-AI peak		% change from pre-AI peak
	Pro se count	Year			Pro se count	Year	
FL - Middle	684	2021	87.4%	OH - Northern	555	2013	35.3%
FL - Northern	235	2020	16.2%	OH - Southern	192	2011	39.6%
FL - Southern	599	2015	65.1%	OK - Eastern	55	2021	-5.45%
GA - Middle	210	2016	34.8%	OK - Northern	65	2010	52.3%
GA - Northern	1030	2012	50.6%	OK - Western	101	2018	88.1%
GA - Southern	152	2011	1.32%	OR	242	2017	32.6%
HI	140	2011	6.43%	PA - Eastern	652	2012	21.2%
IA - Northern	46	2022	100%	PA - Middle	353	2017	15.9%
IA - Southern	63	2012	57.1%	PA - Western	187	2018	81.3%
ID	83	2012	39.8%	RI	99	2009	16.2%
IL - Central	168	2018	25.0%	SC	380	2011	1.84%
IL - Northern	760	2012	62.9%	SD	59	2019	33.9%
IL - Southern	165	2011	-57.6%	TN - Eastern	115	2020	93.9%
IN - Northern	175	2010	57.7%	TN - Middle	187	2017	56.1%
IN - Southern	308	2021	32.1%	TN - Western	239	2011	10.9%
KS	178	2010	20.8%	TX - Eastern	278	2016	55.0%
KY - Eastern	81	2018	46.9%	TX - Northern	540	2013	52.4%
KY - Western	127	2010	40.2%	TX - Southern	529	2011	70.1%
LA - Eastern	452	2018	-53.1%	TX - Western	370	2022	89.5%
LA - Middle	60	2018	80.0%	UT	136	2011	44.1%
LA - Western	779	2022	-66.0%	VA - Eastern	482	2018	76.1%
MA	393	2018	70.0%	VA - Western	116	2018	4.31%
MD	642	2015	32.6%	VT	47	2022	2051%
ME	78	2015	47.4%	WA - Eastern	84	2012	34.5%
MI - Eastern	273	2018	4.40%	WA - Western	497	2011	29.0%
MI - Western	276	2012	40.2%	WI - Eastern	247	2018	22.7%
MN	289	2019	46.4%	WI - Western	135	2013	19.3%
MO - Eastern	245	2014	41.2%	WV - Northern	69	2012	2.90%
MO - Western	302	2017	-31.8%	WV - Southern	80	2009	-17.5%
MS - Northern	64	2010	60.9%	WY	41	2008	-9.76%

B.5. Case Durations

We also report the *pro se* case resolution rate at the $W = 90, 365$ observation windows, as well as companion plots for represented cases at all three horizons. At $W = 365$, the sample does not include any cases filed in 2025.

B.6. Docket-entry burden

Pro se case resolution rate within 90 days of filing, by monthly filing cohort

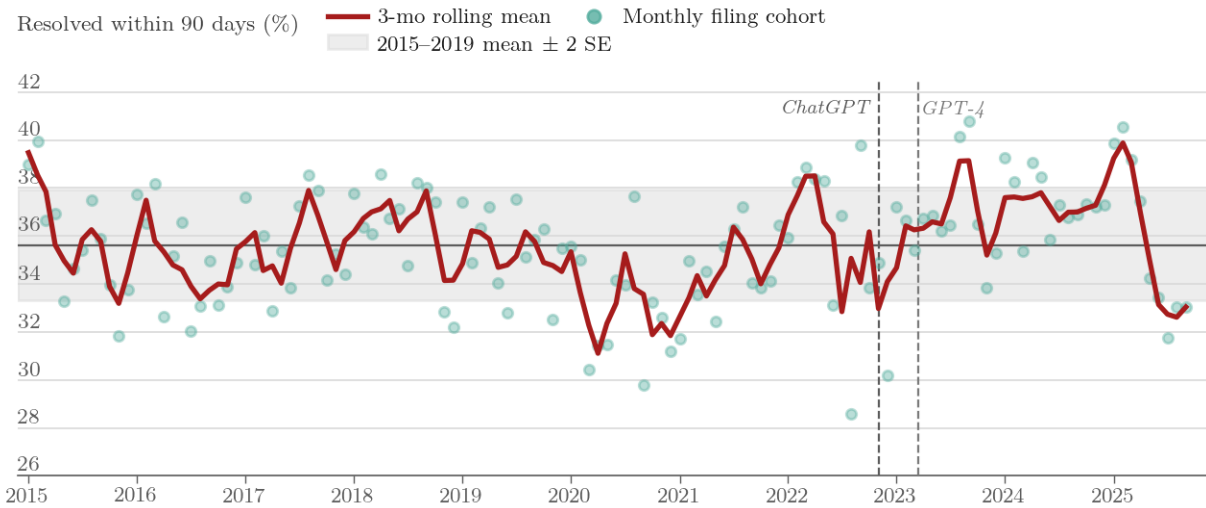


Figure B9. Monthly cohort *pro se* case resolution rate at 90 days from filing. Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

Pro se case resolution rate within 365 days of filing, by monthly filing cohort

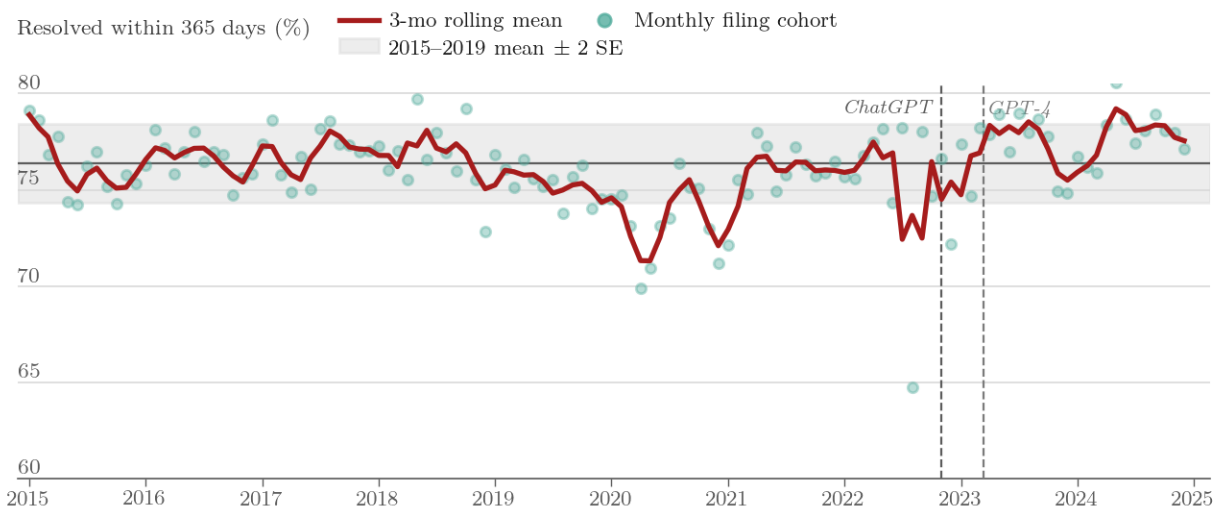


Figure B10. Monthly cohort *pro se* case resolution rate at 365 days from filing. Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

Represented case resolution rate within 90 days of filing, by monthly filing co

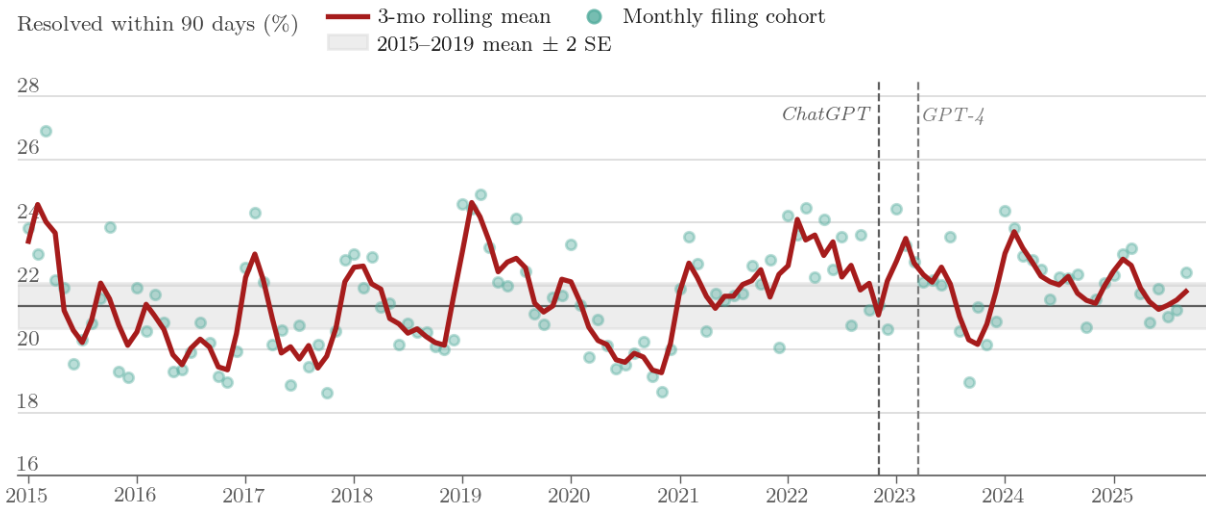


Figure B11. Monthly cohort represented case resolution rate at 90 days from filing. Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

Represented case resolution rate within 180 days of filing, by monthly filing c

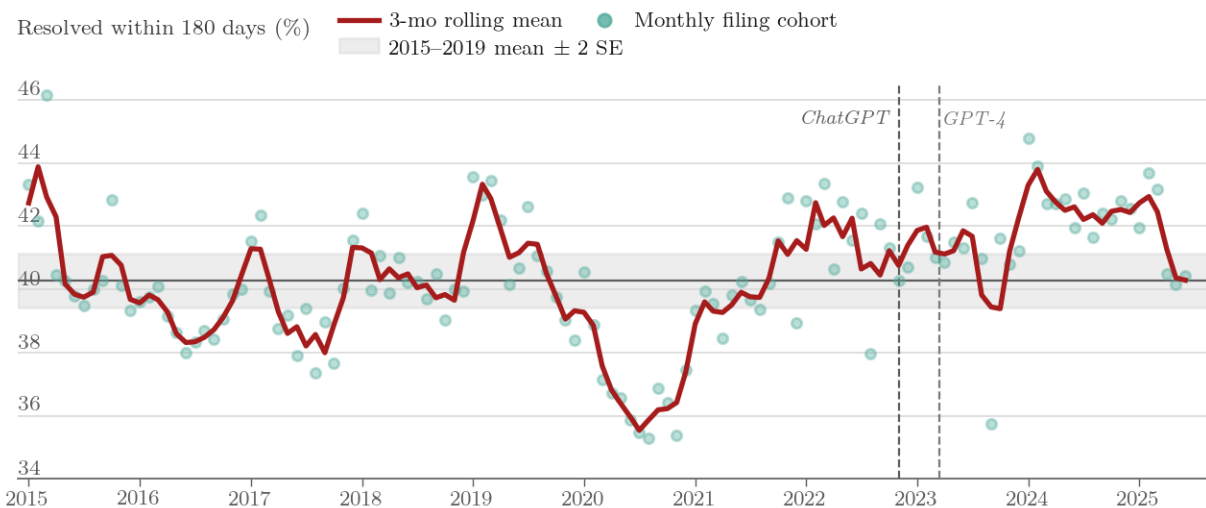


Figure B12. Monthly cohort represented case resolution rate at 180 days from filing. Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

Represented case resolution rate within 365 days of filing, by monthly filing c

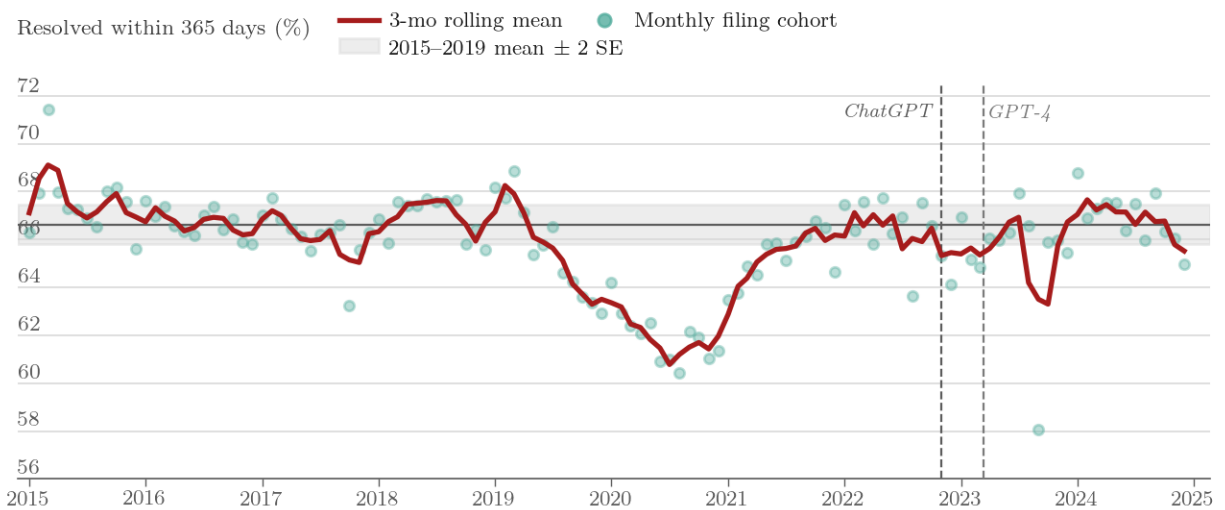


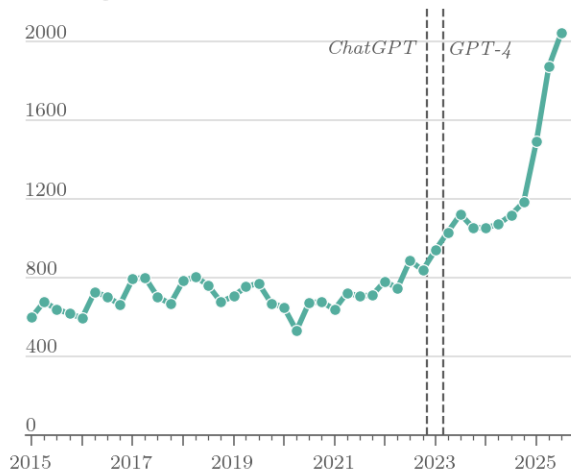
Figure B13. Monthly cohort represented case resolution rate at 365 days from filing. Source: FJC IDB, non-prisoner non-MDL, obs-window filtered.

Civil litigation docket-entry burden, by representation status

Pro se

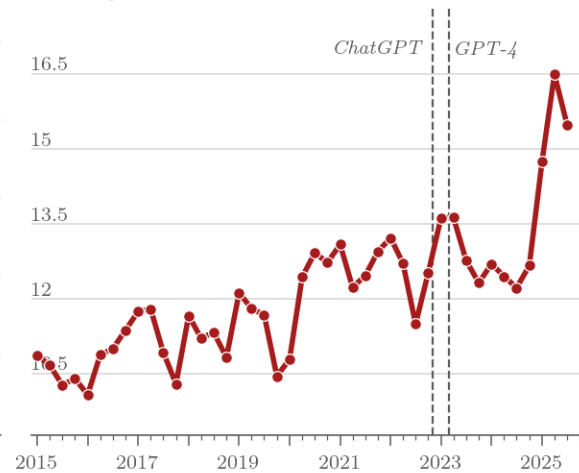
90-day entries per court

Entries per court



90-day entries per case

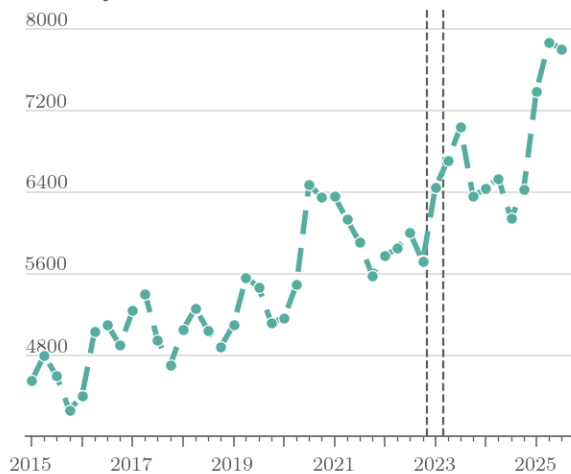
Entries per case



Represented

90-day entries per court

Entries per court



90-day entries per case

Entries per case



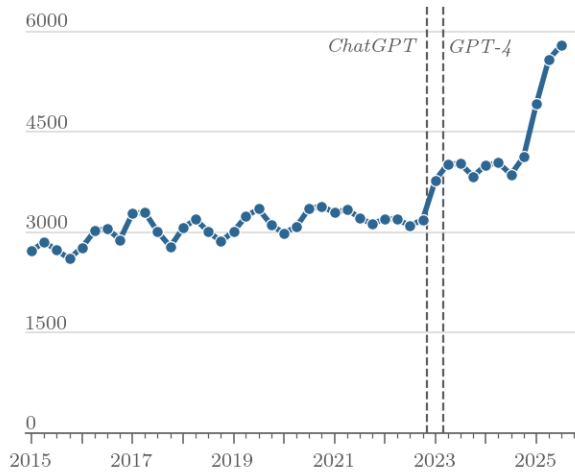
Figure B14. Docket entries in the first 90 days of a case, quarterly, for 54 full-feed districts. Top row: *pro se*; bottom row: represented. Left: entries per court. Right: entries per case. Source: PACER docket metadata; FJC IDB.

Civil litigation docket-entry burden, by case complexity

Simple

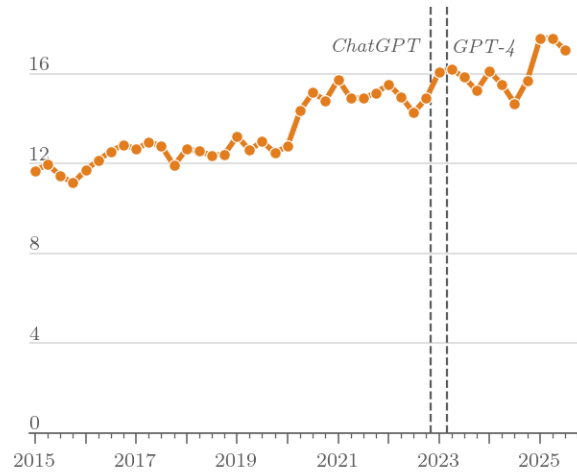
90-day entries per court

Entries per court



90-day entries per case

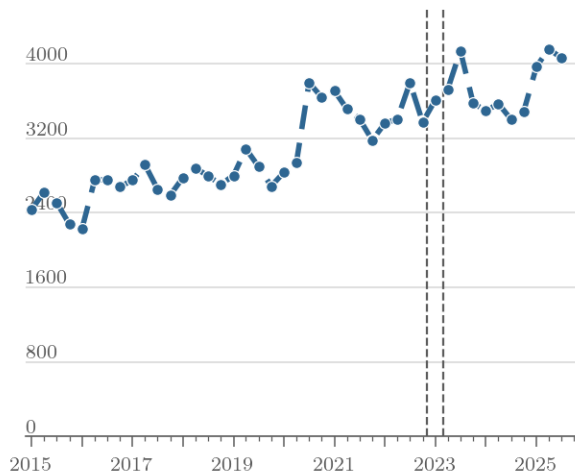
Entries per case



Complex

90-day entries per court

Entries per court



90-day entries per case

Entries per case

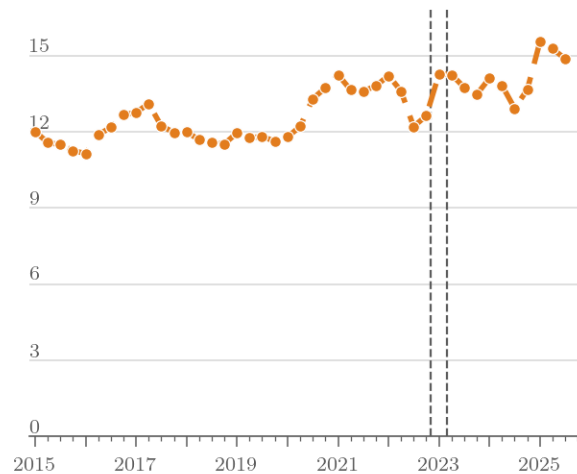
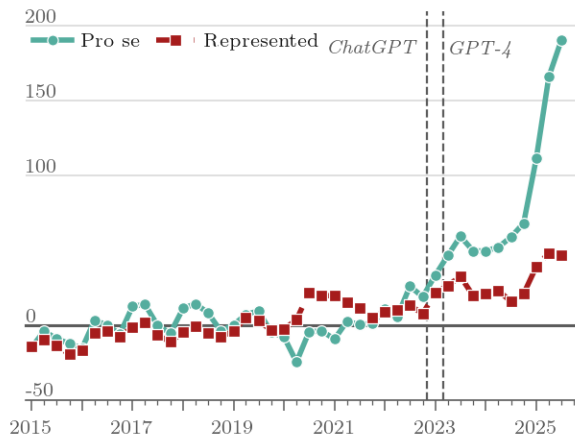


Figure B15. Docket entries in the first 90 days of a case, quarterly, for 54 full-feed districts, split by simple/complex NOS partition. Top row: simple NOS codes. Bottom row: complex NOS codes. Left: entries per court. Right: entries per case. Source: FJC IDB and PACER docket metadata.

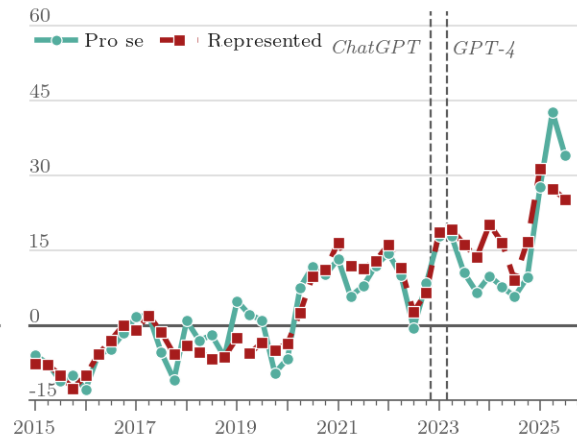
Civil litigation docket-entry burden, relative to pre-AI baseline

Pro se vs represented

90-day entries per court
Change from pre-AI mean (%)

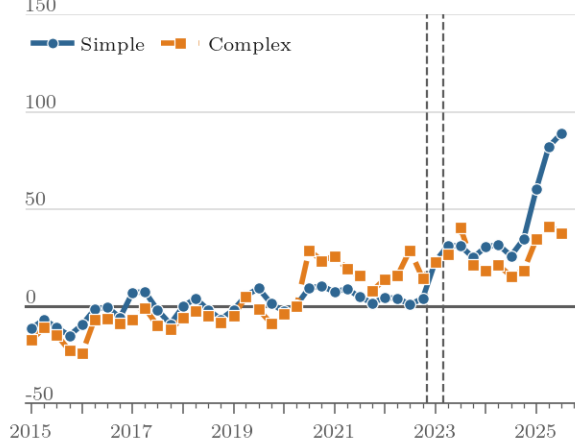


90-day entries per case
Change from pre-AI mean (%)



Simple vs complex

90-day entries per court
Change from pre-AI mean (%)



90-day entries per case
Change from pre-AI mean (%)

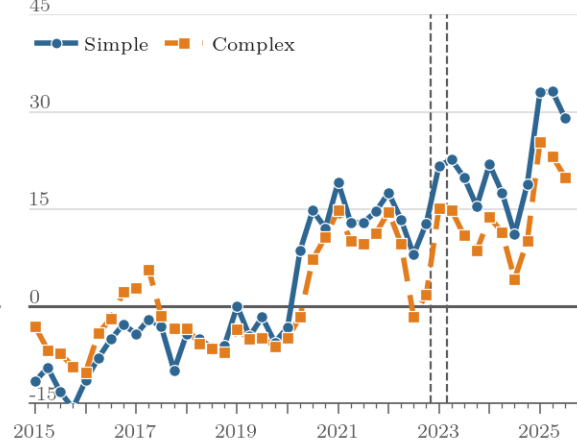


Figure B16. Docket entries in the first 90 days of a case, quarterly, for 54 full-feed districts, shown as percent deviations from the pre-AI mean. Top row: *pro se* vs. represented. Bottom row: simple vs. complex. Left: entries per court. Right: entries per case. Source: FJC IDB and PACER docket metadata.