

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

So much for plain language: An analysis of the accessibility of United States federal laws (1951-2009)

Permalink

<https://escholarship.org/uc/item/22n2b92b>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 44(44)

Authors

Martinez, Eric

Mollica, Francis

Gibson, Edward

Publication Date

2022

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

So much for plain language: An analysis of the accessibility of United States federal laws (1951-2009)

Eric Martínez (ericmart@mit.edu)

Brain & Cognitive Sciences, MIT
Cambridge, USA

Francis Mollica (mollicaf@gmail.com)

School of Informatics, University of Edinburgh
Edinburgh, UK

Edward Gibson (egibson@mit.edu)

Brain & Cognitive Sciences, MIT
Cambridge, USA

Abstract

Over the last 50 years, there have been efforts on behalf of the US government to simplify public legal documents for the benefit of society at large. However, there has been no systematic evaluation of how effective these efforts—collectively referred to as the “plain language movement”—have been. Here we report the results of a large-scale longitudinal corpus analysis ($n \approx 225$ million words), in which we compare every law passed by congress between 1951 to 2009 (as well as concurrent resolutions and proclamations), with a comparably sized sample of English texts from four different genres published during the same time period. We find that laws remain laden with features associated with processing difficulty—including center-embedding, passive voice, low-frequency jargon and capitalization—relative to each of the four baseline genres of English, and that the prevalence of these features has not meaningfully declined since the onset of the plain language movement (in some cases, their prevalence has increased). These findings suggest that top-down efforts to simplify legal language have thus far remained largely ineffectual, despite the apparent tractability of these changes, raising and informing difficult questions of law and public policy.

Keywords: law and language; natural legal language processing; law and cognitive science; psycholinguistics

Introduction

Ignorantia juris non excusat is an ancient maxim of the law which holds that “ignorance of the law is no excuse” (Garner et al., 2004). This ancient maxim remains at the heart of modern legal systems, which typically presume that the public understands the entirety of the legal doctrine and, consequently, do not typically allow ignorance or mistakes of the law as a defence to a crime (Institute, 1984; Arsanjani, 1999). Of course, the presumption that a nation’s citizenry is aware of the content of its laws does not appear to be well-grounded in fact. While part of the public’s ignorance of the law may be attributed to a mere lack of exposure, it seems intuitively obvious that when the public does attempt to understand legal documents they have difficulty doing so. Indeed, the difficulty of reading legal texts has long been acknowledged not just by those tasked with reading these documents but by those creating these documents as well. Sporadic attempts to draw up laws in “simple language, using words that everyone could understand” date back as far back as the eighteenth century in Europe (Mattila, 2016), but have mostly been ignored (Adler, 2012).

In the United States, top-down efforts to simplify government documents for the benefit of the public began as early as the 1970s, when Richard Nixon mandated that the Federal Registry be drafted in “layman’s terms” and Jimmy Carter issued Executive Orders intended to make government regulations “easy-to-understand by those who were required to comply with them” (Exec. Order No. 13648, 1979; Plain Language Action Information Network, 2011). These and subsequent attempts to make government language more accessible have been collectively referred to as the “plain language movement.” The most recent call-to-arms, the Plain Writing Act of 2010, established formal guidelines regarding how to write government documents clearly for a lay audience (Plain Writing Act of 2010, n.d.).

The plain language movement spurred research exploring how to best simplify specific cases of public-facing legal language, such as jury instructions (Charrow & Charrow, 1979; Heuer & Penrod, 1989) and Miranda warnings (Goldstein, Condie, Kalbeitzner, Osman, & Geier, 2003; Rogers, Harrison, Shuman, Sewell, & Hazelwood, 2007). Many of the insights from this literature, as well as the general psycholinguistic literature, are now reflected in the Federal Plain Language Guidelines. While these studies have successfully demonstrated that replacing problematic features of legal text (such as archaic legal jargon and complex syntax) with “plain English” equivalents increases comprehension rates among laypeople, they apply only to a small portion of the total corpus of legal language and are less relevant to people’s experience with the legal system than actual laws.¹ However, there remains no systematic analysis of to what extent the plain-language movement impacted the accessibility of federal laws. Moreover, on a more general level, there also remains no systematic evaluation of the accessibility of federal laws over time relative to more standard forms of English.

¹For example, although jury instructions can be an important part of cases that go to trial, a small and diminishing percentage of civil and criminal cases actually go to trial (as low as 3% for the former and 5% for the latter: (Refo, 2004; Rakoff, Daumier, & Case, 2014)). Moreover, while Miranda warnings provide crucial information to criminal suspects in police custody, the majority of individuals’ contact with legal language takes place outside the context of criminal or civil suits.

To address these questions, we conducted a corpus analysis of (a) every law passed by congress between 1951 and 2009 (as well as concurrent resolutions not signed into law and proclamations issued by the president), and (b) a large sample of magazine articles, newspaper articles, non-fiction books and fiction books published over the same time span. We analyzed a variety of linguistic and stylistic features, whose use is (a) discouraged by the Federal Plain Language Guidelines, (b) associated with language processing difficulty in psycholinguistic research, and (c) purportedly common in legal documents. We find that each of these features remains strikingly more prevalent in public legal documents relative to each of our baseline texts, with virtually none having significantly decreased in prevalence since the start of the plain language movement.

Materials and Methods

Corpus Materials

For our analysis we constructed an exhaustive corpus of every public law, private law, concurrent resolution and proclamation issued by the American federal government between the years 1951 and 2009 using publicly available online resources from the United States library of congress (Library of Congress, 2021). As a baseline, we extracted a comparably-sized sample of English texts drawn from the Corpus of Historical American English (Davies, 2012), which consisted of a broad sample of fiction books, non-fiction books, magazine articles, and newspaper articles also published between 1951 and 2009.

To process and analyze these corpora, we used a number of natural language processing tools. One of the primary tools we used was the Stanford Stanza natural language package (Qi, Zhang, Zhang, Bolton, & Manning, 2020), a state-of-the-art NLP toolkit which we used to tokenize each document into sentences, lemmatize and tag each word by part of speech, and syntactically parse each tokenized sentence. Stanza has been shown to achieve over 90% accuracy on a variety of NLP tasks (Qi et al., 2020). To verify its accuracy on our specific corpora and for our specific metrics, we spot-checked a random sample of 1000 sentences across our corpora by (a) hand-coding whether a given sentence had a passive-voice structure or a center-embedded clause, and (b) for each sentence comparing whether the parser's judgments aligned with the hand-coded judgments. Using this method, we found that the parser was 97.93% accurate at detecting by-passive structures (95% CI: 97.04 to 98.82) and 88.95% accurate at detecting center-embedding structures (95% CI: 86.98 to 90.73).

We also used the SUBTLEX word frequency dictionary (Brysbaert & New, 2009) to get a word frequency estimate as a proxy for how common a given word in each corpus appears in everyday speech. The SUBTLEX frequency values themselves are derived from a large-scale corpus of American film subtitles and have been shown to correlate with reading-time behavior (Brysbaert & New, 2009). We also used WordNet

(Miller, 1995), which, in tandem with SUBTLEX, was used to estimate whether a given word could have been replaced by a higher frequency word with the same meaning.

Pre-processing for both corpora were identical. Sentences were first tokenized and dependency-parsed using the Stanford Stanza NLP package. We then removed sentences without punctuation, as well as those with fewer than 10 words so as to remove headings, which are not really sentences but would otherwise be counted as such. We also removed sentences with 3+ consecutive punctuation marks so as to get rid of more non-sentences in both corpora. The total number of words after filtering was 225,899,179 (68,031,729 words for the legal corpus and 157,867,450 for the non-legal corpus). After filtering out non-sentences, we then dependency-parsed each corpus, lemmatized and tagged each word by part of speech and computed our indices of processing difficulty.

Indices of Processing Difficulty

In each of these corpora we sought to determine the prevalence of six features that are associated with processing difficulty in the general psycholinguistic literature, whose use is discouraged by the Federal plain language guidelines, and which are purportedly common in legal documents Fig. 1. Below is a description of each feature, as well as our method of computing it in each corpus.

Word frequency. For each of our corpora we sought to determine, on average, how frequently the words in said corpora occur in everyday speech. Words that are infrequently used in everyday speech cause comprehension difficulties for readers relative to higher-frequency synonyms (Marks, Doctorow, & Wittrock, 1974). Legal language is reportedly laden with low-frequency jargon, such as *aforesaid*, *hereinafter*, and *to wit* (Rayner, Ashby, Pollatsek, & Reichle, 2004), and recent work has shown the language in contracts to be lower-frequency than that of other genres of English (Martinez, Mollica, & Gibson, 2021). According to the official plain-language guidelines, government writing should avoid the use of such low-frequency “dry legalisms” and “jargon” (Plain Language Action Information Network, 2011).

Frequency values were extracted from the SUBTLEX corpus of American film subtitles (Brysbaert & New, 2009). To avoid including non-content words, we limited our analysis of frequency to the words in our corpora marked as a verb, noun, adjective or adverb according to Stanza. Proper nouns and other words that did not appear in the SUBTLEX corpus received a score of NA.

Word choice. Although many argue that the processing difficulty of unfamiliar language is a necessary consequence of the specialized concepts and corresponding terminology used to refer to those concepts by lawyers (cf. (Tobia, 2020)), recent work suggests that private legal documents contain a high-proportion of overly complicated language that can be replaced with simpler terms that have the same meaning (Martinez et al., 2021). The official plain-language guidelines encourage the use of “familiar or commonly used” words over such “unusual,” “obscure” or “unnecessarily complicated lan-

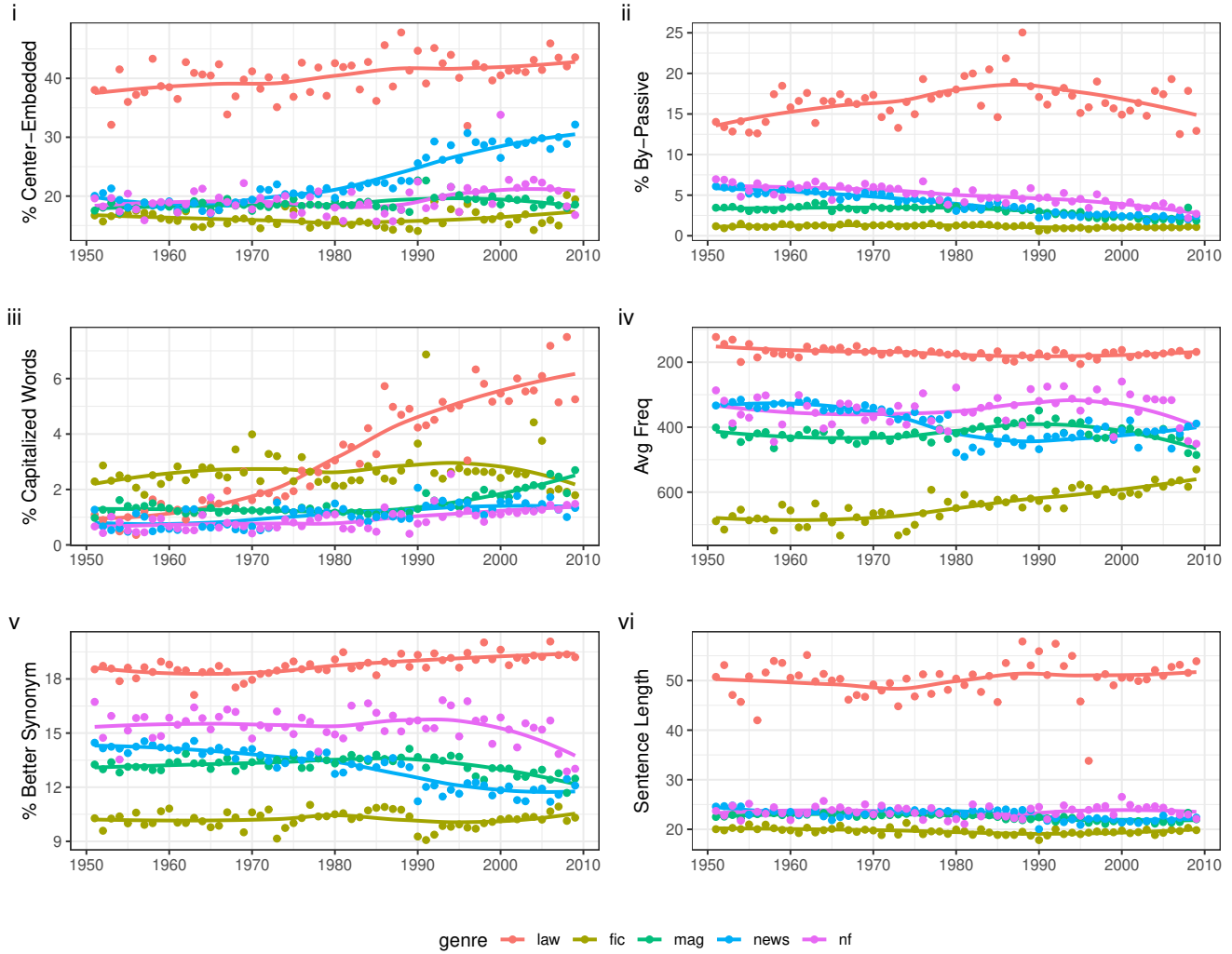


Figure 1: Comparison of indices of linguistic processing difficulty in federal laws vs four genres of standard English, including fiction books, magazine articles, newspaper articles, and non-fiction books (1951-2009).

guage” (Plain Language Action Information Network, 2011). Here we sought to quantify the amount of unnecessarily complicated language in federal laws by calculating the percentage of words in each corpus that could have been replaced with a higher-frequency synonym.

We conducted two versions of this analysis using two separate assumptions. First, we made the conservative assumption that the authors intended the least common sense of each word used in a given corpus because while legal terms may resemble common words in form, they may have a more specialized meaning, such as the concept of “consideration” in contract law (American Law Institute and National Conference of Commissioners on Uniform State Laws, 2002). Second, we made an anti-conservative assumption that the authors intended the most common sense of each word in a given corpus. Again, we limit our analysis to verbs, common nouns, adjectives and adverbs. For both analyses, we determined the least (conservative) or most (anti-conservative) common meaning/sense of that word according to WordNet

(Miller, 1995). For all words sharing that meaning/sense (i.e., synonyms), we looked up the SUBTLEX frequency value and coded whether the SUBTLEX frequency value of any synonym was higher than that of the actual word used in the text (1=Yes; 0=No). Results of the conservative and anti-conservative method did not differ. Therefore, we report the conservative method.

Capitalization. In each corpus we computed the percentage of words that contained non-standard capitalization (specifically, those that were in ALL CAPS). Although the plain-language guidelines do not discourage the use of all-capitalization in government writing, evidence suggests that non-standard capitalization (“ALL WARRANTIES ARE HEREBY DISCLAIMED”) is common in certain types of private legal documents (Martinez et al., 2021) and has shown to inhibit comprehension in older readers (Arbel & Toler, 2020), relative to standard capitalization. In our analysis we coded a word as being in “all-caps” by calculating the proportion of alphabetic word tokens that were marked by Stanza as being

entirely in uppercase letters.

Center-embedded clauses. Plain-language guidelines discourage the use of “convoluted” sentences, particularly those that are “loaded with dependent clauses” and which separate the “essential parts” of a sentence from each other (i.e. the subject, verb and object). The most notorious examples of such sentences contain center-embedded structures, in which a sentence or clause is embedded within the center of another sentence or clause (“*all such payments and benefits, including the payments and benefits under Section 3(a) hereof, being hereinafter referred to as the ‘Total Payments’*”). Center-embedded structures cause processing difficulty for readers (Gibson, 1998) and have been shown to inhibit recall of legal content relative to clauses that have been un-embedded into separate sentences (Martinez et al., 2021). Here we calculated the percentage of sentences in each corpus containing a center-embedded clause. We coded a sentence as containing a center-embedded clause if a predicate dependent clause as parsed by Stanza (i.e. clausal subjects, clausal complements, open clausal complements, adjectival clauses, and adverbial clauses) was followed by a word as opposed to an end-of-sentence punctuation mark.

Sentence Length Plain-language guidelines encourage the use of shorter sentences so as to “break the information up into smaller, easier-to-process units” (Plain Language Action Information Network, 2011). Legal texts, especially laws and other public documents, are reportedly filled with long sentences (Hiltunen, 2012). Although evidence suggests sentence length is less of a predictor of processing difficulty than center-embedding and other types of syntactic complexity (Marton & Schwartz, 2003), to err on the side of caution we include sentence length in our analysis. We computed sentence length by calculating the number of words in each sentence as determined by Stanza.

Passive-voice structures. Federal Plain Language Guidelines advocate for using the active voice instead of the passive voice. Passive-voice structures are acquired later than active-voice structures and have been shown to pose comprehension difficulties for adults in certain circumstances, particularly in the context of implausible sentences e.g. “*the girl was kicked by the ball*” (Ferreira, 2003). Although (Martinez et al., 2021) recently found evidence that passive voice structures did not inhibit recall of legal content relative to active-voice structures in contracts, it may be that the stimuli used in (Martinez et al., 2021) did not span the circumstances shown to induce the comprehension errors seen in adult experiments. To err with caution, we include passive voice structures in our analysis, particularly reversible passives or *by*-passives (e.g. “*the information shall be maintained by the Federal Government*” as opposed to “*the information shall be maintained*”), which can be more easily replaced by active-voice structures without a loss or distortion in meaning. We coded a sentence as containing a reversible passive voice structure if a word was marked with the passive voice features by Stanza and had the word *by* in the same head according to the Stanza parse.

Results

Efficacy of the Plain Language Movement

Were the plain-language movement to have been effective, one would expect (a) the prevalence of difficult-to-process features to have meaningfully decreased over time, and (b) the decrease to coincide with the onset of the plain-language movement. To evaluate this prediction, for each of our six indices of processing difficulty we conducted a break-point Bayesian regression limited to the legal-corpus data. The break point was fixed at 1972, a plausible year for the plain language movement’s call-to-arms. If the plain language movement was effective, one would expect the slope of the regression line for after the breakpoint (i.e., 1972-2009) to be both negative and less than the slope of the regression line before the breakpoint (i.e. 1951-1972). For word frequency and sentence length, we used a linear regression to predict the mean value of these metrics per sentence. For all other indices, we used binomial logistic regression.²

Regression coefficients for all indices can be found in Table 1. Our models revealed the plain language movement to have coincided with no meaningful decrease in any of our indices.

General Trends in Accessibility

Even if the plain language movement did not coincide with a decrease in difficulty-inducing structures in legal texts, it may be the case that (a) difficulty-inducing structures became more prevalent in other texts relative to or as well as legal language, or that (b) legal language was not filled with very high indices of difficulty-inducing structures to begin with. To evaluate these alternative accounts, as well as to obtain a more general systematic account of the accessibility of federal laws—both temporally and relative to other genres of English—we first computed the descriptive statistics of each of index within the corpora over time. We found that for each year, the prevalence of virtually every metric was higher in federal laws than in any of the four genres of the plain-language corpus (in most cases, the difference was striking). Visualizations of these results can be viewed in Figure 1.

We then used Bayesian regression methods to estimate the influence of corpus (legal vs baseline) over time (in years) for each of our indices of processing difficulty. For every metric, our models revealed federal laws to contain more difficult to process structures than our baseline texts. For every metric except capitalization, our models revealed no meaningful influence of time on the prevalence of a given metric, nor of the interaction between time and corpus. For each index, we first considered two models: one with a main effect of Corpus and Year, and one with an additional interaction term for Corpus and Year. A Bayes-factor comparison for each index except center-embedding revealed at least moderate evidence

²In the case of our sentence-level metrics (center-embedding and passive voice), the regression estimated the influence of our predictors on whether a sentence had a given metric. In the case of our word-level metrics (capitalization and word choice), the regression estimated the influence of our predictors on whether a word had a given metric.

	Intercept	Before 1972	After 1972
Word Frequency	-1.87 [-1.87 - -1.87]	0.02 [0.02-0.02]	0.03 [0.03-0.03]
Word Choice	-1.28 [-1.28 - -1.28]	0.00 [0.00-0.00]	0.00 [-0.00-0.00]
Capitalization	-3.59 [-3.59 - -3.58]	0.03 [0.03-0.03]	0.05 [0.05-0.05]
Center-embedding	-0.38 [-0.39 - -0.38]	0.00 [0.00-0.00]	0.00 [0.00-0.00]
Sentence Length	-2.10 [-2.10 - -2.10]	0.3 [0.02-0.02]	0.03 [0.03-0.03]
Passive Voice	-1.19 [-1.20 - -1.18]	-0.01 [-0.01 - -0.01]	0.02 [0.02-0.02]

Table 1: Estimates and 95% confidence intervals for the intercept and slopes of the breakpoint regression models.

(BF 100) for the second model over the first model. We therefore only report the results of the second model in Table 2

Discussion

As discussed above, the present study had two main aims. The first aim was to investigate to what extent federal laws have grown more accessible since the start of the plain language movement. The fact that for each of our regression models the slope of the line after 1972 was either positive or greater than the slope of the line before 1972 suggests that laws have not gotten meaningfully simpler since the onset of the plain-language movement.

With regard to the second aim, we next investigated to what extent federal laws deviate from so-called “plain English.” As visualized and documented above, all of the metrics we looked at which were associated with psycholinguistic complexity (i.e. “non-plain English”) were startlingly more prevalent in federal laws than each of our baseline texts, not just overall but for virtually every year between 1951 and 2009. In line with common intuition and plain-language advocates and consistent with recent findings regarding private legal documents (Martinez et al., 2021), this suggests that public legal language deviates quite heavily from plain English, and has been and continues to be more difficult to understand than standard English.

Our study provides the first systematic large-scale account of the accessibility of public legal language—both longitudinally and compared to more standard forms of English—and provides an even starker account of the efficacy of plain-language efforts than previously assumed. Whereas current plain-language advocates describe progress as “way slow” and acknowledge that “much remains to be done to improve” (Plain Language Action Information Network, 2011), our results instead suggest that despite the movement’s best efforts, progress may in-fact be non-existent, at least with regard to laws, resolutions and proclamations prior to 2010.

Having documented the profile of public legal language over the last 50 years and demonstrated the inefficacy of plain-language efforts over the same time period, further extensions to this study—with regard to academic scholarship and government advocacy—should seek not only to confirm the extent to which these findings hold for more recent laws and other types of government documents, but also to understand the cause of the complexity of legal language. In other words, not only how lawyers and lawmakers write but

why they choose to write they way that they do.

One possibility is that the style in which laws are currently written is necessary to maintain communicative precision. This possibility is undercut by our results, which focused on features that are known to have simpler alternatives (e.g. “this law prohibits smoking in public areas” versus “smoking in public areas is prohibited by this law”), as well as previous findings that show comprehension of legal content with a simplified register (e.g., Masson & Waldron, 1994; Martinez et al., 2021). While it seems entirely plausible that certain legal jargon is inevitable, our results suggest that in many instances such jargon can be replaced with simpler alternatives that preserve meaning.

Another possibility is that esoteric text arises out a mismatch between the priorities of the writer and reader of a law. If lawmakers’ priorities differ from the reader’s priorities they may even do this implicitly as opposed to engaging in an outright “conspiracy of gobbledegook” (Mellinkoff, 2004). Lastly, lawmakers may not *choose* to write in an esoteric manner. Similar to the “curse of knowledge” (Hinds, 1999; Nickerson, 1999), they may not realize that their language is too complicated for the average reader to understand (Azuelos-Atias, 2018). If true, this would predict that the processing difficulty of legal texts may be alleviated as lawmakers become more aware of both the ways in which public legal documents tend to be complex, as well as the alternatives available to them in order to make them less complex. Further work into the plausibility of these hypotheses could yield insight into how best to persuade lawmakers to integrate the findings of our and similar studies and help alleviate the mismatch between the ubiquity and impenetrability of legal texts in the modern era.

References

- Adler, M. (2012). The plain language movement. In *The oxford handbook of language and law*.
- American Law Institute and National Conference of Commissioners on Uniform State Laws. (2002). *Uniform commercial code (U.C.C.) s 2-216(2)*.
- Arbel, Y. A., & Toler, A. (2020). All-caps. *Journal of Empirical Legal Studies*, 17(4), 862–896.
- Arsanjani, M. H. (1999). The rome statute of the international criminal court. *American Journal of International Law*, 93(1), 22–43.

	Intercept	Corpus	Year	Corpus:Year	BF
Word Freq.	344.68 [341.30-348.13]	171.07 [167.44-175.11]	-0.32 [-0.56- -0.08]	-0.52 [-0.74- -0.29]	Inf
Word Choice	0.20 [0.09-0.30]	-4.73 [-4.92-4.54]	-0.00 [-0.00-0.00]	0.00 [0.00-0.00]	Inf
ALL CAPS	-3.88 [-3.88 - -3.88]	-0.13 [-0.13- -0.13]	0.02 [0.02-0.02]	-0.01 [-0.01- -0.01]	Inf
Embedding	-0.97 [-0.97- -0.97]	-0.57 [-0.57- -0.56]	0.00 [0.00- 0.00]	-0.00 [-0.00- -0.00]	12.13
Sent. Length	35.99 [35.60-36.30]	-14.73 [-15.03- -14.34]	.02 [-0.00-0.04]	-0.04 [-0.06- -0.02]	Inf
Passive Voice	18.34 [17.70-18.97]	-22.56 [-23.53 -21.67]	-0.01 [-0.01- -0.01]	0.01 [0.01-0.01]	Inf

Table 2: Estimates and 95% confidence intervals for the intercept and slopes of the Bayesian regression models, as well as the Bayes Factor (BF) estimates in favor of these models over models with an interaction term.

- Azuélos-Atias, S. (2018). Making legal language clear to legal laypersons. *Legal Pragmatics*, 288, 101.
- Brysbaert, M., & New, B. (2009). Moving beyond kučera and francis: A critical evaluation of current word frequency norms and the introduction of a new and improved word frequency measure for american english. *Behavior research methods*, 41(4), 977–990.
- Charrow, R. P., & Charrow, V. R. (1979). Making legal language understandable: A psycholinguistic study of jury instructions. *Columbia law review*, 79(7), 1306–1374.
- Davies, M. (2012). Expanding horizons in historical linguistics with the 400-million word corpus of historical american english. *Corpora*, 7(2), 121–157.
- Exec. Order No. 13648. (1979). *44 fr 69609*.
- Ferreira, F. (2003). The misinterpretation of noncanonical sentences. *Cognitive psychology*, 47(2), 164–203.
- Garner, B. A., et al. (2004). Black's law dictionary.
- Gibson, E. (1998). Linguistic complexity: Locality of syntactic dependencies. *Cognition*, 68(1), 1–76.
- Goldstein, N. E. S., Condie, L. O., Kalbeitzer, R., Osman, D., & Geier, J. L. (2003). Juvenile offenders' miranda rights comprehension and self-reported likelihood of offering false confessions. *Assessment*, 10(4), 359–369.
- Heuer, L., & Penrod, S. D. (1989). Instructing jurors. *Law and Human Behavior*, 13(4), 409–430.
- Hiltunen, R. (2012). The grammar and structure of legal texts. In *The oxford handbook of language and law*.
- Hinds, P. J. (1999). The curse of expertise: The effects of expertise and debiasing methods on prediction of novice performance. *Journal of Experimental Psychology: Applied*, 5(2), 205.
- Institute, A. L. (1984). *Model penal code*.
- Library of Congress. (2021). *United states statutes at large*.
- Marks, C. B., Doctorow, M. J., & Wittrock, M. C. (1974). Word frequency and reading comprehension1. *The Journal of Educational Research*, 67(6), 259–262.
- Martinez, E., Mollica, F., & Gibson, E. (2021). *Long-distance syntactic dependencies drive the complexity of legal language*. Retrieved from <https://doi.org/10.31234/osf.io/hfbdt>
- Marton, K., & Schwartz, R. G. (2003). Working memory capacity and language processes in children with specific language impairment.
- Masson, M. E., & Waldron, M. A. (1994). Comprehension of legal contracts by non-experts: Effectiveness of plain language redrafting. *Applied Cognitive Psychology*, 8(1), 67–85.
- Mattila, H. E. (2016). *Comparative legal linguistics: language of law, latin and modern lingua francas*. Routledge.
- Mellinkoff, D. (2004). *The language of the law*. Wipf and Stock Publishers.
- Miller, G. A. (1995). Wordnet: a lexical database for english. *Communications of the ACM*, 38(11), 39–41.
- Nickerson, R. S. (1999). How we know—and sometimes misjudge—what others know: Imputing one's own knowledge to others. *Psychological bulletin*, 125(6), 737.
- Plain Language Action Information Network. (2011). *Federal plain language guidelines*.
- Plain Writing Act of 2010. (n.d.). "5 U.S.C. § 301. 2010".
- Qi, P., Zhang, Y., Zhang, Y., Bolton, J., & Manning, C. D. (2020). Stanza: A python natural language processing toolkit for many human languages. *arXiv preprint arXiv:2003.07082*, 0.
- Rakoff, J. S., Daumier, H., & Case, A. C. (2014). Why innocent people plead guilty. *The New York Review of Books*, 20.
- Rayner, K., Ashby, J., Pollatsek, A., & Reichle, E. D. (2004). The effects of frequency and predictability on eye fixations in reading: Implications for the ez reader model. *Journal of Experimental Psychology: Human Perception and Performance*, 30(4), 720.
- Refo, P. L. (2004). The vanishing trial. *Journal of Empirical Legal Studies*, 1(3), v–vii.
- Rogers, R., Harrison, K. S., Shuman, D. W., Sewell, K. W., & Hazelwood, L. L. (2007). An analysis of miranda warnings and waivers: Comprehension and coverage. *Law and*

human behavior, 31(2), 177–192.

Tobia, K. P. (2020). Legal concepts and legal expertise. *Social Science Research Network*, 0.