

Annual Review of Law and Social Science

Computational Methods in Legal Analysis

Jens Frankenreiter¹ and Michael A. Livermore²

¹Ira M. Millstein Center for Global Markets and Corporate Ownership, Columbia Law School, Columbia University, New York, NY 10027, USA

²School of Law, University of Virginia, Charlottesville, Virginia 22903, USA;
email: mlivermore@virginia.edu

Annu. Rev. Law Soc. Sci. 2020. 16:39–57

First published as a Review in Advance on
July 17, 2020

The *Annual Review of Law and Social Science* is online
at lawsocsci.annualreviews.org

<https://doi.org/10.1146/annurev-lawsocsci-052720-121843>

Copyright © 2020 by Annual Reviews.
All rights reserved

**ANNUAL
REVIEWS CONNECT**

www.annualreviews.org

- Download figures
- Navigate cited references
- Keyword search
- Explore related articles
- Share via email or social media

Keywords

computational law, machine learning, empirical legal studies, legal history, digital humanities, big data

Abstract

The digitization of legal texts and advances in artificial intelligence, natural language processing, text mining, network analysis, and machine learning have led to new forms of legal analysis by lawyers and law scholars. This article provides an overview of how computational methods are affecting research across the varied landscape of legal scholarship, from the interpretation of legal texts to the quantitative estimation of causal factors that shape the law. As computational tools continue to penetrate legal scholarship, they allow scholars to gain traction on traditional research questions and may engender entirely new research programs. Already, computational methods have facilitated important contributions in a diverse array of law-related research areas. As these tools continue to advance, and law scholars become more familiar with their potential applications, the impact of computational methods is likely to continue to grow.

INTRODUCTION

In recent decades, the digitization of legal texts and advances in information processing technology and theory have worked to transform the practice and study of law. Techniques from the fields of artificial intelligence, natural language processing, text mining, network analysis, and machine learning are now routinely taken up by legal practitioners and law scholars. These new tools are being applied both to practical challenges that arise in the lives of lawyers and legal subjects and to scholarly research questions concerning law and legal institutions.

This article provides an overview of the many exciting new uses of computational methods in legal analysis, with an emphasis on the use of these techniques by scholars. Traditionally, the core research methodology of law scholars involved reading and interpreting legal texts, such as statutes or judicial opinions, and then making descriptive, predictive, or normative claims about the law and legal decision making. This mode of analysis serves as the basis of classical doctrinal analysis, interdisciplinary work in humanistic fields such as legal history and jurisprudence, and law-related social science work. Although this method has supported centuries of scholarship, it is also limiting, in terms of both the amount of data that can be processed and the types of analysis that are possible.

New computational technologies open important new research opportunities for law scholars by expanding the analytic methods that can be applied to legal texts. One such opportunity, which has been pursued for several decades by scholars in law and computer science, studies the law by formally representing legal rules directly as executable code. More recently, a group of scholars have focused on translating legal texts directly into data that can be subject to quantitative analysis. This law-as-data approach uses computer-based tools to extract useful information from high-dimensional legal data sets, and in particular from collections of legal documents. This information can be analyzed to gain traction on long-standing research questions within law scholarship and can also engender new research programs that were beyond the grasp of traditional methods.

Law-as-data techniques have very general applicability for law scholarship, which is a diverse world of research that encompasses a wide range of questions, disciplines, and approaches. One classic distinction is between internal and external perspectives on law. Internal law scholarship participates in legal discourse by taking seriously the reasons and norms offered by legal actors; traditional doctrinal analysis falls into this category. External scholarship, by contrast, focuses on analyzing the behavior of legal actors as a social phenomenon; historical or social science research is often characterized as external scholarship. As we discuss below, law-as-data techniques have contributed to both internal and external scholarship on a wide range of questions.

Other distinctions in law scholarship arise out of the different research goals that can be pursued, which include summarization and synthesis of legal content; interpretation of the law within its social or political context; normative analysis of the desirability of legal standards or practices; prediction of the behavior of legal actors; and causal identification of the factors that affect legal change or the influence of law on other social phenomena. To the extent that these research projects involve the analysis of legal texts (or other textual sources), computational tools can usefully be brought to bear.

The structure of this review is as follows. We first distinguish two general approaches to applying computational method to the law: the law-as-code approach and the law-as-data approach. The first attempts to model the law as a set of formal rules, whereas the latter uses computational techniques to extract information from legal texts to incorporate into more general research programs. We focus on law as data in this review because it touches nearly the entire landscape of law scholarship. We then briefly introduce techniques that are used for extracting and summarizing quantitative information from legal texts and then discuss some of the many applications where

these law-as-data tools have been put to use by researchers. We conclude with a discussion of some of the practical implications of these developments and current challenges for this budding field.

LAW AS CODE VERSUS LAW AS DATA

In recent decades, a group of legal academics and computer scientists have sought to develop computer code that represents legal rules contained in statutes and case law (Ashley 1990, Bench-Capon 1991, Branting 1991, Gardner 1987, Rissland & Skalak 1991, Sergot et al. 1986). Such knowledge representation approaches conceive of the law as a set of logical rules that can ultimately be formalized as inputs that a computer can process directly. A simple example of a legal knowledge representation system would be a decision tree with a series of binary choices that results in a determination of liability under a tort regime.

In practice, some legal knowledge representation systems have had enormous success; for example, the field of tax preparation has been transformed by such software (Contos et al. 2011). Legal commentators have discussed the possibility that future lawmakers might shift from the familiar natural language of statutes and regulations to a form of executable code in a common programmable language that would be clearer and potentially tailorable to individual contexts (Casey & Niblett 2017, Coglianese 2004, Fagan & Levmore 2019).

An alternative approach to applying computational methods to the law focuses on the potential to treat the law as data. With the explosion of digitized textual resources and advances in computational natural language processing, social scientists have begun to explore text as a source of data on phenomena of interest, a trend that has been described as using text as data (Grimmer & Stewart 2013). Within the humanities, a similar trend has been referred to as distant reading, an approach to understanding texts based on statistical analysis rather than traditional close reading (Moretti 2013). Text as data and distant reading have made important inroads in their respective disciplines, leading to new empirical results and theoretical innovations, and in recent years, researchers have begun to explore their application to the law (Livermore & Rockmore 2019).

Law as data sits at the intersection of quantitative empirical legal studies and traditional, doctrinally oriented scholarship. From the perspective of quantitative empirical legal studies, law as data can be, simply, more grist for the mill. The same techniques of statistical analysis that have been used in the field for decades can be applied to textual data as well as any other data (although with some important complications and caveats, discussed below). Textual data also facilitate new quantitative techniques that are related to but distinct from prior work in empirical legal studies. At the same time, law as data shares with the domain of traditional scholarship its emphasis on the special place of legal texts. Although the quantitative tools that can be used to understand law as data may be less familiar, the questions and interests of traditional legal scholarship are often amenable (at least theoretically) to investigation using law-as-data approaches.

Evans et al. (2007) provide an early example of law as data through an empirical lens. For that project, the authors apply a word-frequency model to *amicus curiae* briefs submitted to the Supreme Court to generate classifications along an ideological dimension. At the opposite end of the internal/external spectrum are efforts to promote the use by courts of corpus linguistics as a tool of legal interpretation, and in particular as a means to discern the original meaning of legal texts (Lee & Mouritsen 2018, Strang 2016, Tobia 2020). In research aims and disciplinary tools, the social science approach taken by Evans et al. (2007) is quite remote from the doctrinally grounded intent of Lee & Mouritsen (2018) and other proponents of legal corpus linguistics—what they share is the conversion of legal texts to data, which they then turn to their various uses via the analytic techniques of their disciplines.

Although both research areas draw heavily from computational concepts, they are quite different in their goals and applications. Law as code is a quite specific research program of its own, whereas law as data represents a suite of general-purpose analytic tools. These tools can shed light on a diverse set of questions across the landscape of law scholarship, both building on and expanding prior research programs. The balance of this review focuses on law-as-data research. Bench-Capon et al. (2012) and Rissland et al. (2003) provide useful reviews of law-as-code research, which is an important paradigm that has practical implications.

FROM LAW TO DATA

The defining feature of law-as-data tools is the application of quantitative, mathematical techniques to legal texts. Traditionally, legal texts have been analyzed qualitatively or, more recently, have been converted to data for purposes of quantitative analysis manually through the process of hand coding. As a practical matter, the time and effort involved in reading individual documents (for either qualitative analysis or hand coding) limits the scale of a corpus that can be analyzed. In addition, both qualitative interpretation and hand coding are subjective, and the latter represents documents in terms of a small number of binary characterizations, such as whether a legal issue was present, leading to potentially useful information being discarded.

More sophisticated computational techniques that can convert legal texts to data have many advantages over traditional techniques but raise other challenges. Translating the natural language of the law to numerical values involves a range of choices and trade-offs that are best made in light of the research question at issue. In this section, we introduce some common approaches for quantitatively representing legal language, along with some of their advantages and drawbacks. Subsequent sections show these techniques in use to help illustrate both their value and their limitations in different research contexts.

The Curse of Dimensionality

Legal documents are natively very high-dimensional objects. For example, if texts are treated as an ordered sequence of words, the dimensionality of such a representation would be extremely high: Documents that were 1,000 words long in a simple language with a vocabulary of 1,000 words would be represented in a space of $10^{3,000}$ dimensions. Such high dimensionality implies that, without a vast number of observations, almost all of the possible outcomes in the language space of 1,000 words by 1,000 positions would be unobserved, in large part because they would form sequences that violated the syntax of the language or were semantically nonsensical. This mismatch between the number of variables and the amount of data makes quantitative analysis difficult.

The goal for any method of quantitatively representing textual documents is to achieve dimensionality reduction in a way that preserves information while compacting the data into a smaller (i.e., lower-dimensional) space. One approach to carrying out this task is to create quantitative variables by hand, by reading documents and assigning numerical values based on predefined categories. Computational tools allow for other options.

Relatively Simple Metrics

Different approaches to representing text as quantitative information vary substantially in terms of their complexity. Maybe the easiest and most straightforward way to transform text into data is to focus on simple statistics summarizing basic characteristics of the texts under investigation.

Examples of this approach include studies measuring the length of judicial opinions (Black & Spriggs 2008) and comparably simple measures of writing style based on features such as word or sentence length, in particular readability scores (Feldman 2019, Law & Zaring 2010, Potter 2019, Whalen 2015).

Another example of relatively straightforward measures that can be extracted from legal texts is based on references to other documents (i.e., citations). One particularly convenient way to do so is the use of a network representation. Under this approach, legal documents are represented as rows and columns in a matrix in which each entry indicates whether a document cites another document (Fowler & Jeon 2008, Lupu & Fowler 2013). Potentially, such information is augmented based on whether a citation is positive or negative (Clark & Lauderdale 2010, Cross 2012) or based on the number of times another document is mentioned (Clark & Lauderdale 2012). Network representations of legal documents are not limited to judicial opinions: Statutory cross-references are another context in which network representations have been deployed (Badawi & Dari-Mattiacci 2019, Katz & Bommarito 2014).

Richer Representation of Substance and Style

Researchers in fields such as natural language processing and computational text analysis have developed tools to provide more information-dense representations of documents. Law scholars have begun to apply these more sophisticated tools to legal texts such as judicial opinions and statutes.

Three basic sets of building blocks are used to create more sophisticated representations of legal texts. The most common group of approaches is based on the bag-of-words construct, which represents documents as term-frequency vectors—i.e., lists of words in a document and their frequencies (e.g., Roberts et al. 2016). Because the order of words in a document is ignored in a bag-of-words representation, the dimensions of a term-frequency vector representation of a document are equal to the vocabulary in a corpus. Term-frequency vectors can also be generated using longer sequences, so-called *n*-grams, where *n* is the length of sequences. A 2-gram model would represent documents as term-frequency vectors in which each term is two words long.

Researchers face several further technical choices when converting texts into term-frequency vectors. For example, analysts must decide whether to treat inflected forms of the same word as different words (e.g., whether to stem or lemmatize words), how to treat proper nouns and numbers, and whether to use a weighting schedule such as the term frequency-inverse document frequency (tf-idf) scaling scheme. These choices can have consequences for the information that is ultimately analyzed and must be made carefully (Grimmer & Stewart 2013). It is also possible to focus on specific categories of words (rather than all words in the vocabulary) based on the research question at hand. For example, Carlson et al. (2016) conduct a stylometric analysis of US Supreme Court decisions based on term-frequency vectors of only function words, which are nonsemantic words (such as *if* and *or*). Frankenreiter (2019) and Langford et al. (2020) extend this approach to the European Court of Justice and investment arbitration awards, respectively. The Linguistic Inquiry and Word Count tool offers various ways to summarize documents based on word frequency, for instance, whether there is a high occurrence of intensifiers such as the word *clearly*, and has been used to study language in court documents (Black et al. 2016). Sentiment analysis focuses on words that convey emotional valence and has been used to study judicial opinions as well as public comments received by administrative agencies (Busch & Pelc 2019, Carlson et al. 2016, Livermore et al. 2018, Rice & Zorn 2019).

Most law-as-data research has, to date, been based on bag-of-words models. However, in natural language, the order of words matters, and so a second set of building blocks can be constructed

that take account of word order. This can be done by taking advantage of expert knowledge about language to parse text based on semantic structure (Ashley & Brüninghaus 2009). Such an approach can capture relationships that are lost in bag-of-words models, such as whether an adjective is used to modify one or another noun. Hand coding the semantic structure of language is laborious to do at scale, but some automated parsers have been developed.¹ One challenge of this approach is that it risks exploding the dimensionality of a data representation.

Another way to capture word order that is less subject to the dimension explosion problem is via algorithms that represent words as vectors in a multidimensional space based on the proximity to other words in a document or corpus. One such approach is the skip gram representation that serves as the basis for word-embedding models of language and related document embeddings approaches (Mikolov et al. 2013). Recent work has extended the embeddings model in the legal context to the level of opinions and judges (Ash & Chen 2019, Rice et al. 2019).

The third set of building blocks are more exotic and application specific, inviting creativity on the part of the researcher based on the question at hand. One example is the use of plagiarism-detection software, which seeks to identify unusual similarity in relatively long sequences of words. Such software has been used to track the influence of lower court opinions and parties' briefs on Supreme Court opinions (Corley 2008, Corley et al. 2011). In a similar vein, Choi & Gulati (2005) rely on compression scores to compare the authorship of judicial opinions.

Downscaling

Depending on the techniques employed, the ensuing textual representation can be rather high dimensional, potentially necessitating an additional step of dimensionality reduction. Two main families of machine-learning approaches are commonly used to achieve this goal: supervised and unsupervised. Under supervised approaches, a researcher uses a training set of labeled data in combination with a machine-learning model to predict certain features in out-of-sample data. Unsupervised approaches do not require labeled data but instead amount to sophisticated pattern-sensitive quantitative summaries of texts.

Rauterberg & Talley (2017) provide an example of a supervised approach. In this paper, the authors examine a large corpus of corporate documents to determine the pervasiveness of provisions governing the freedom of corporate officers to pursue outside business opportunities. To identify these provisions, the authors select a sample set of documents, hand code the relevant provisions, and then use the sample as a training set for a machine-learning classifier. Once an acceptably low error rate is achieved, the authors use the classifier to identify similar provisions in the remaining documents and then conduct further quantitative analyses based on these data. Rauterberg & Talley's (2017) supervised approach radically reduces the dimensionality of the documents down to a single binary prediction of whether or not it contained a "corporate opportunity waiver."

An example of an unstructured approach to representing legal documents is via a topic model, which seeks to identify latent subject matter categories in unstructured corpora of texts (Blei 2012, Blei & Lafferty 2007). The "topics" that are discovered by a topic model are distributions over the vocabulary in the corpus. Documents are described as distributions over those topics, thus preserving some of their semantic content while reducing the number of dimensions to the number of topics (typically between 10 and 100). Topic models have been used in several papers to model legal documents (Carter et al. 2016, Corley et al. 2011, Livermore et al. 2017). The outputs of topic models can be difficult to interpret; model selection, parameterization, and validation are

¹ Both Google and the Natural Language Processing Group at Stanford have publicly available parsers (Petrov 2016, Stanford Nat. Lang. Proc. Group n.d.).

active areas of research (Caspi & Stiglitz 2020). Simpler examples of unsupervised approaches to dimensionality reduction include principal component analysis and comparable techniques.

Like many unstructured representations of legal language, network representations of legal documents are potentially also rather high dimensional. In these situations, network analysis provides a range of additional possibilities to downscale the complexity of the data. For example, studies can use various measures of network centrality as a measure of the importance of a document within an interlinked corpus (Fowler & Jeon 2008, Fowler et al. 2007, Lupu & Fowler 2013, Lupu & Voeten 2012) or tools such as spectral clustering to examine the relatedness of documents (Carlson et al. 2016, Frankenreiter 2019).

APPLICATIONS USING TEXTUAL DATA

As mentioned before, law-as-data techniques can be applied in a broad variety of settings. The following overview provides several examples of how researchers have used these tools to pursue different types of research questions using a range of techniques. Given the growing breadth of law-as-data research, the following discussion is meant to be illustrative rather than exhaustive.

Causal Inference

We start our overview with applications of law as data in studies that focus on quantitatively investigating causal relationships. Work in this field uses explicit causal models, study design, and parameter estimation to identify causal relationships between different variables using tools such as regression analysis.

In principle, the law can take on different roles in this kind of work. In research investigating the consequences of legal arrangements on real-world outcomes, numerical representations of the law typically feature as independent variables in a regression. The majority of existing studies using computational methods, however, are concerned with shedding light on factors influencing the legal process itself. In this context, legal variables are used mostly as dependent variables.

Although regression analysis and similar techniques have been a centerpiece of much of the existing quantitative empirical legal research, these methods are less well suited to dealing with the high-dimensional data made available by using text as data. In recent years, empirical researchers in a range of fields have begun discussing how empirical research can overcome this challenge by leveraging the power of machine learning, which, in principle, is well equipped for processing of high-dimensional data.

There are a range of different proposals on how machine learning can be used to address the challenges of text-related dimensionality in causal inference research. Two broad categories can be identified. First, machine-learning techniques can be used to assist in the creation of variables from textual data, which are subsequently used in the analysis just like any other variable. Second, the credibility of regression analysis and similar techniques can be improved by using machine-learning techniques in certain steps of the analysis (Copus et al. 2019, Mullainathan & Spiess 2017). Our survey of the literature reveals that existing work in the legal field generally belongs to the first category. Work in the second category may increase in the future; for example, recent research in neighboring disciplines showcases how textual data can be leveraged to tackle inferential challenges pertaining to causal identification (e.g., Roberts et al. 2020).

In current applications, causal research using text as data thus essentially adopts a two-step approach. In the first step, textual information is condensed down to a small number of variables, resulting in a drastic reduction of dimensionality. In the second step, these variables are used in standard statistical examinations, such as regression analysis. Because of the dimensionality

reduction achieved in the first step, the data create no major challenges for traditional statistical tools. Researchers use many different computational tools to generate low-dimensional variables, and differences in researchers' objectives will also often imply the use of different techniques.

In some cases, researchers are simply interested in investigating the determinants of certain low-dimensional textual features of opinions. Examples include studies of the factors influencing the length and clarity of judicial opinions (Black & Spriggs 2008, Black & Wedeking 2016, Goelzhauser & Cann 2014, Leonard & Ross 2016, Whalen 2015) and standard-form contracts (Marotta-Wurgler & Taylor 2013). In a similar vein, Hinkle et al. (2012) study the use of hedging and intensifying language in opinions by district court judges depending on how their own ideology compares with that of judges at the courts of appeals. Finally, Wahlbeck et al. (2002) use similar metrics to compute a measure for the stylistic differences between all opinions by Justice Marshall and Justice Powell in the 1985 term, which they use to analyze differences in the style of opinions assigned to different law clerks.

In a second set of studies, researchers use computational methods to translate textual data into proxies for other characteristics of legal documents. Particularly relevant in this context are attempts to measure the “legal” content of judicial opinions or other legal texts. One way to do so is to use computational methods to extend the scale of an analysis that would otherwise rely on hand-coded information. In this context, researchers will typically use supervised machine-learning algorithms to replicate an existing coding of a subset of cases for the entire data set. Alongside Rauterberg & Talley (2017) (discussed above), examples in this category include work by Talley & O’Kane (2012), who investigate force majeure provisions in merger and acquisition contracts, and Nyarko (2019), who uses a similar approach based on supervised machine learning to determine whether contracts include choice-of-forum provisions. Rice (2014) studies the effect of Supreme Court decisions on lower courts by training a range of supervised algorithms to predict issue categories in the text of opinions based on a prior hand-coded data set. Alschner (2017), in an examination of the impact of investment arbitration on investment protection treaties, uses a coding procedure based on identifying keywords to identify whether investment protection treaties contain certain provisions.

Researchers can also use computational methods to obtain measures for the characteristics of legal texts that cannot be obtained by way of human coding. This approach is most often used in the context of legal opinions. Rice (2017) develops a method of topic concentration based on a standard topic model, which he then uses to investigate whether dissenting opinions at the Supreme Court force the majority to cover a broader set of issues in their opinion. Carlson et al. (2020) use a topic model to proxy for the substantive legal issues present in cases before US appellate courts to test for interactions between judicial characteristics and the types of cases that lead to published opinions. In an effort to establish the extent to which parties can influence the content of Supreme Court decisions, Corley (2008) uses plagiarism software to establish the degree to which the language in these opinions mirrors that in the parties’ briefs. Using the same approach, Corley et al. (2011) investigate the influence of lower court opinions on the writing of Supreme Court justices. Oldfather et al. (2012) use a slightly different technique to examine the similarity between party briefs and the wording of a set of opinions of the First Circuit. Their analysis relies on comparing the frequency distributions of words as well as the citations in the relevant documents. Finally, Carlson et al. (2016) use the frequency distribution of function words to construct measures for the consistency of the writing style of the Supreme Court and its justices, which they use to gain insights into the changing role of clerks at the court.

Similar approaches have also been used to measure the contents of other legal texts. Stiglitz (2014) studies “midnight rulemaking” by administrative agencies, using a measure for the level of controversy associated with administrative rules based on the text in the rule and preamble. In

another study, Stiglitz (2018) examines whether differences in the nondelegation doctrine at the state level translate into differences in the delegation of rulemaking authority to administrative agencies in different states. For this, he extracts information on the extent and quality of such delegations from new laws adopted in the states. Alschner & Skougarevskiy (2016) investigate the consistency of a country's investment treaties based on a measure for the similarity of the language used in these agreements. Kosnik (2014) develops different measures for the "completeness" of contracts in a study analyzing the determinants of the flexibility of agreements. Beuve et al. (2019) and Moszoro et al. (2016) investigate differences in the rigidity of public and private contracts. Finally, McLane (2019) examines the benefits and costs of the use of boilerplate language with various measures of the amount of boilerplate in securities disclosures.

Finally, information in legal documents can also serve as a proxy for other facts that cannot be observed directly. For example, Smith (2014) uses the frequency distribution of specific words in circuit court opinions to determine the degree to which the case leading to the opinion hinges on factual and/or legal issues; this measure is then used to investigate whether judge ideology plays out differently in different cases. Patton & Smith (2017) measure the impact of attorney gender on the frequency with which attorneys are interrupted by judges during oral arguments. As a basis for their analysis, they extract data from transcripts of Supreme Court oral arguments to establish proxies for the length of uninterrupted speech by each attorney delivering an argument.

Prediction and Classification

Existing studies within the causal inference paradigm do not use high-dimensional textual data to investigate the research question directly. Instead, they condense textual information into one or a small number of variables and then carry out more familiar statistical analyses based on the transformed, lower-dimensional data. This process necessarily results in a loss of information. An alternative approach is to retain the high-dimensional data but use different analytic tools. In particular, researchers have begun to take advantage of machine-learning algorithms, which are relatively better equipped to deal with high-dimensional data compared with traditional statistical techniques such as regression analysis.

This shift does come with a cost: Machine learning is primarily geared toward prediction and classification rather than causal inference (Mullainathan & Spiess 2017). To take advantage of these tools, quantitative researchers attempt to identify noncausal research questions that can be answered using machine learning.² In these studies, the prediction or classification task constitutes the main element of the analysis.

One application of machine-learning tools is to categorize legal texts by assigning labels to (parts of) legal documents based on textual patterns. Early examples include attempts to train algorithms to identify whether certain legally relevant fact patterns are present in a case (Brüninghaus & Ashley 1997) and to locate legally relevant text in opinions (Daniels & Rissland 1997). Later work sought to assign documents to subject matter categories (Gonçalves & Quaresma 2005, Thompson 2001), to predict the authoring judge from the vocabulary used in an opinion (Li et al. 2013, Rosenthal & Yoon 2011), and to generate classifications of legal texts along an ideological dimension (Evans et al. 2007).³ Dumas (2019) tests whether it is possible to predict the political affiliation of judges based on the language in judicial opinions, and Hausladen et al. (2020) train a machine-learning algorithm to predict the ideological direction of case outcomes from

²More generally, Kleinberg et al. (2015) refer to these problems as "prediction policy problems."

³Clark (2017) discusses an unpublished paper by McGuire & Vanberg (2005) that set out to measure the ideological leanings of Supreme Court opinions based on the frequency distribution of words used in the opinion texts.

the text of US appellate court opinions. Recent work has used these methods to apply labels to contracts (Lippi et al. 2017) and privacy policies (Contissa et al. 2018). Although these papers explore purely predictive questions, the data from this type of prediction exercise can be used in studies that are interested in uncovering causal relationships, as in work by Rauterberg & Talley (2017).

Another application is to predict the outcomes of legal cases. Katz et al. (2017) build on work by Ruger et al. (2004)—which predicted the outcomes of Supreme Court cases in a single term based on a simple statistical model using a small set of variables—by expanding the number of variables, using a more sophisticated prediction algorithm, and predicting outcomes of the court’s cases over most of its existence. The first attempt to predict the outcomes of cases based on text appears to be by Ashley & Brüninghaus (2009). This paper predicts the outcomes of cases indirectly by first estimating the presence of certain fact patterns in the case. Aletras et al. (2016) show that it is possible to predict legal outcomes directly from text with relatively high accuracy. However, this study bases its analysis on the texts of the judgments that contain the outcomes that the model sets out to predict. Alexander et al. (2019) report early efforts to extract information from docket sheets to use early-stage information (such as the identity of parties and their lawyers) to predict litigation outcomes in employment law cases.

Interpretation and Description

The conversion of law into data that can be quantitatively analyzed and summarized creates an altogether new scale for more traditionally qualitative disciplines, such as legal history. One important relatively early work in this vein is that of Klingenstein et al. (2014), which examines transcripts of criminal cases tried in London during the early modern period, finding that the testimony in violent cases became more distinctive over the period from the late 1700s to the early twentieth century, tracking evolving social attitudes toward violence during that time.

There are several other applications of law as data in historical work. Romney (2016) constructs term-frequency representations of judicial opinions collected in the first eight volumes of the *Hawaii Reports* to examine how jurists’ interpretation of the interaction of racial status and access to the writ of habeas changed over the course of the mid-to-late nineteenth century. Nystrom & Tanenhaus (2016a,b) construct and analyze a data set of state session laws to examine the adoption of harsh laws affecting juvenile criminal defendants in the 1990s. Young (2013) constructs a topic model of newspaper articles published during the 18-year period after the end of the Civil War to examine whether the citizenry was more engaged in questions of constitutional design at that time. Funk & Mullen (2018) investigate the migration of the Field Code through the American South and West. Each of these projects integrates quantitative text analysis tools with methods and questions familiar to legal historians.

Quantitative text analysis has also been used to examine temporal change more broadly. Rockmore et al. (2017) examine the diffusion of constitutional concepts over space and time. Livermore et al. (2017) combine a topic model and a machine-learning classifier to examine whether the content of Supreme Court opinions has become less court-like over the course of the twentieth century, as compared with a baseline of opinions in the federal appellate courts. Carlson et al. (2016) examine temporal trends in Supreme Court sentiment, finding a trend toward more negative language over time. Other works that assess temporally related stylistic trends on the court include those by Johnson (2014) and Feldman (2017). Cross & Pennebaker (2014) focus on opinions issued during the Roberts court to examine stylistic and linguistic features of the court’s writings during that time period. Even more tightly focused, Varsava (2018) examines the writings of Justice Gorsuch when he was a judge on the Tenth Circuit prior to his elevation

to the Supreme Court, to examine how his stylistic drift might help predict the types of writing style choices he might make in the future. In a similar vein, Ash & Chen (2018) use different measures of Justice Kavanaugh's past decisions to make inferences about his ideological leanings. Pozen et al. (2019) investigate the polarization of constitutional speech in the US Congress over time.

Another important application of law-as-data tools is in developing new descriptive metrics that can be used as the basis for quantitative analysis or in the service of interpretive tasks, such as synthesizing legal doctrine. Descriptive analyses of this sort can help develop and refine theory and also facilitate the measurement of variables of interest, which is essential to quantitative causal and predictive analysis.

For example, several researchers have applied quantitative tools to describe legal complexity. Qualitatively, legal complexity could be understood in the terms laid out in the Federalist No. 62 (James Madison), as "laws [] so voluminous that they cannot be read, or so incoherent that they cannot be understood." Researchers have developed various quantitative metrics that attempt to track this qualitative notion. Owens & Wedeking (2011) use linguistic indicators from the Linguistic Inquiry and Word Count tool to estimate the cognitive complexity of the language in US Supreme Court opinions. Ruhl & Katz (2015) and Ruhl et al. (2017) draw from the field of complexity theory to propose quantitative measures of legal complexity, and Katz & Bommarito (2014) use network characteristics (including hierarchical structure and cross-references) and linguistic characteristics akin to readability scores to measure the relative complexity of titles in the US Code. Bommarito & Katz (2017) use text from Securities and Exchange Commission filings to identify references to US regulations or agencies as an estimate of the complexity of the regulatory environment.

Topic models have frequently been used to provide high-level compact representations of the semantic content of texts in a corpus, which can be used as the inputs into quantitative analysis or can be subject to qualitative description or interpretation. Quinn et al. (2010) apply a topic model to text of speeches in the Congressional Record to extract issue categories that captured congressional attention in the period from 1997 to 2004. Rice (2019) carries out a similar exercise for the US Supreme Court, arguing that the more fine-grained representations topic models provide can be superior to the dichotomous indicators of opinion content often used in empirical study of judicial opinions. Lauderdale & Clark (2014) use results from a topic model alongside voting data to create a more fine-grained ideological description of US Supreme Court justices. Law (2016) and Ruhl et al. (2018) examine how well the output of unstructured models maps onto expert-generated categories of constitutional preambles and executive pronouncements, respectively. Law (2018) uses topic models to examine human rights language across international instruments and constitutions and argues that there are two general dialects within a broader shared language community.

Researchers can also use law-as-data metrics to delve into questions that are defined by a legal area rather than a corpus. For example, Fagan (2015) and Macey & Mitts (2014) use topic models in this way, the first to inform a taxonomy of how US courts use successor liability to achieve a set of legal outcomes, and the second to examine the rationales used by courts in veil-piercing cases. Rice et al. (2019) use an alternative tool, word embeddings, to measure the use of racially biased language in judicial opinions.

The study of the precedential force of judicial opinions constitutes yet another area in which computational tools can complement more traditional legal analysis in a meaningful way. The next section discusses studies that use citation data to characterize the treatment of authority in judicial opinions and in causal and predictive models of judicial behavior.

APPLICATIONS USING CITATION DATA

Most work using citation data is concerned with quantitatively investigating the treatment of authority in judicial opinions. Just like with work using textual data, the complexity of the measures varies. Some studies use citation counts and comparable simple statistics, for example, in examinations of the precedential value of decisions. Cross & Spriggs (2010) attempt to identify the most influential Supreme Court opinions and analyze factors predicting whether a case will be cited in the future. Black & Spriggs (2013) focus on factors that influence the depreciation of the citation value of court opinions. Other examples include works by Whalen (2013) and Whalen et al. (2017) that examine network features to predict the likelihood of future citations. Hitt (2016) demonstrates how different features of underlying citation data (in particular, whether the data distinguish between citations that indicate that precedent is followed or not) can influence the outcome of such studies. Others use similar measures to investigate whether the political background of judges affects the set of opinions they cite in their own decisions. Choi & Gulati (2008a) and Niblett & Yoon (2015) address these questions in the context of the Federal Courts of Appeals. Frankenreiter (2017) extends this approach to the European Court of Justice. Choi & Gulati (2008b) test whether judges' citation behavior is influenced not only by their own political background but also by that of fellow panel members. Using a somewhat different approach, Niblett (2010) investigates whether judges at California appeals courts cherry-pick precedents to support their desired outcome.

Other studies use network-analysis techniques to create variables that are then used to tackle similar questions. For example, Lupu & Fowler (2013) use citation data to construct a network-based measure for the embeddedness of Supreme Court opinions in precedent. This measure is then used to test whether the citations in these opinions are influenced not only by case characteristics but also by strategic interaction between the justices. Lupu & Voeten (2012) undertake a similar analysis in the context of the European Court of Human Rights, and Larsson et al. (2017) do so in the context of the European Court of Justice. Fowler et al. (2007) and Fowler & Jeon (2008) develop network-based metrics to generate authority scores that conform to expert judgment about the most important Supreme Court precedents. These scores are used to test several hypotheses about judicial behavior, including the likelihood of overruling a prior decision. One particularly interesting finding, which Carmichael et al. (2017) later confirmed, is that network-based metrics can successfully predict the probability of a future citation of a case.

In addition to using citation data to generate variables for use in traditional statistical examinations, network measures can simply provide a means of summarizing information about citation practices. Early work applies common measures for network analysis to judicial citations (Bommarito et al. 2009, Smith 2007). Clark & Lauderdale (2012) introduce a more legally motivated means of describing dependencies between opinions by building a genealogical model from citation data. Miller (2019) combines data on citations with semantic analysis to examine the development of legal doctrine in Supreme Court intellectual property cases. Related work has investigated the citation networks of courts such as the European Court of Justice (Darlén & Lindholm 2013, 2017; Pelc 2014). Several studies on international courts combine network analysis with a qualitative mode of interpretation (Alschner & Charlotin 2018, Olsen & Küçüksu 2017, Sadl & Olsen 2017, Tarissan & Nollez-Goldbach 2016).

Although the bulk of research using citation data is concerned with judges' citation practices, several studies use these data as proxies for other (unobservable) features of judicial opinions. In particular, such data have been used to generate measures for the ideological leanings of judicial opinions. Clark & Lauderdale (2010), for example, scale Supreme Court opinions in search-and-seizure as well as freedom-of-religion cases based on whether later cases positively or negatively

cite prior cases. This scaling exercise shows substantial overlap with codings of the outcome of the cases, and the authors go on to test various theories about interactions among members of the Supreme Court. Cross (2012) follows a different approach for measuring the ideological leanings of Supreme Court opinions based on the outcomes of later opinions that positively or negatively cite these earlier opinions.

Citations and textual information need not always be separated for purposes of analysis. For example, Leibon et al. (2018) construct network representations of topic model–based semantic information alongside citations in Supreme Court opinions to construct a multi-network model of law search. The authors use this model to predict citation behavior based on the ease of navigation between different documents in this network landscape. Dadgostari et al. (2020) build on this multi-network representation of the law to construct a reinforcement learning model that uses textual information in opinions to predict citations.

PRACTICAL IMPLICATIONS AND CHALLENGES

The work described above is primarily scholarly, oriented to understanding and explaining the law and legal institutions. However, methods and technologies developed for academic uses often turn out to have practical applications as well. This has proven to be the case for both law-as-code and law-as-data techniques. With respect to the former, the basic premise that legal rules can be translated into executable code has found a variety of practical uses, in contexts as diverse as tax preparation (Contos et al. 2011) and family law disputes (Schmitz 2019). In both cases, entrepreneurs have developed knowledge representation systems that reduce at least a portion of the law to what amounts to decision trees that can be navigated by nonexperts, leading to the automation of at least some tasks that were previously performed by legal professionals. A similar law-as-code mindset underlies the application of blockchain technology to contracts (so-called smart contracts) (Raskin 2017) and computer-assisted drafting of legal documents (Betts & Jaep 2017).

Law-as-data techniques have also found their way into legal practice (Remus & Levy 2017). In discovery practice, use of forms of computational text analysis and machine learning is now common, at least in the context of large and complex cases. Commercial tools are now also available to predict the outcomes of at least some types of legal proceedings. Law search is also another area in which there is the potential for broad application of law-as-data techniques, including machine learning and network analysis (Dadgostari et al. 2020, Leibon et al. 2018). Several companies now offer what is advertised as artificial intelligence–assisted law search, although the proprietary nature of such applications makes them difficult to evaluate.

Despite the growing use of computational methods in both academic and practical legal applications, important challenges remain. The disciplinary training of many legal scholars is often limited to the traditional methods of the profession, such as advocacy or doctrinal analysis. The growth of interdisciplinary legal scholarship during the latter part of the twentieth century (especially in the United States) expanded the background of some scholars to include disciplines such as economics, history, and philosophy. However, law scholars with backgrounds in fields such as mathematics, computer science, software engineering, and data science are quite rare. Accordingly, law scholarship and the law school curriculum have—at least arguably—not kept pace with technological developments.

An additional potential set of challenges to the use of computational methods in legal analysis comes from political and legal steps taken by incumbent actors to protect their current positions. Investment in legal start-up companies may be hampered by fears that their services could be understood as “unlicensed practice of law” and therefore subject to prosecution (Hadfield 2008).

Incumbent actors in the United States have also moved aggressively to use copyright law to inhibit the ability of not-for-profit organizations to share access to legal resources, although not always successfully.⁴ Perhaps the most severe step taken to inhibit the field came from the French legislature, which passed a law making it illegal to analyze the decision making of judges using publicly available information that includes their identity.⁵

Notwithstanding these challenges, computational methods are quickly becoming part of the standard analytic tool kit available to scholars of the law as well as legal practitioners. So long as access to legal data continues, and the technologies of natural language processing, machine learning, and computational text analysis improve, it seems likely that these methods will only further work their way into the life of the law.

DISCLOSURE STATEMENT

The authors are not aware of any affiliations, memberships, funding, or financial holdings that might be perceived as affecting the objectivity of this review.

LITERATURE CITED

- Aletras N, Tsarapatsanis D, Preoȃiuc-Pietro D, Lampos V. 2016. Predicting judicial decisions of the European Court of Human Rights: a natural language processing perspective. *PeerJ Comput. Sci.* 2:e93
- Alexander CS, al Jadda K, Feizollahi MJ, Tucker AM. 2019. Using text analytics to predict litigation outcomes. See Livermore & Rockmore 2019, pp. 275–311
- Alschner W. 2017. The impact of investment arbitration on investment treaty design: myths versus reality. *Yale J. Int. Law* 42:1–66
- Alschner W, Charlotin D. 2018. The growing complexity of the International Court of Justice’s self-citation network. *Eur. J. Int. Law* 29:83–112
- Alschner W, Skougarevskiy D. 2016. Mapping the universe of international investment agreements. *J. Int. Econ. Law* 19:561–88
- Ash E, Chen D. 2018. What kind of judge is Brett Kavanaugh? A quantitative analysis. *Cardozo Law Rev. De Novo* 2018:70–100
- Ash E, Chen D. 2019. Case vectors: spatial representations of the law using document embeddings. See Livermore & Rockmore 2019, pp. 313–37
- Ashley KD. 1990. *Modeling Legal Argument: Reasoning with Cases and Hypotheticals*. Cambridge, MA: MIT Press
- Ashley KD, Brüninghaus S. 2009. Automatically classifying case texts and predicting outcomes. *Artif. Intell. Law* 17:125–65
- Badawi AB, Dari-Mattiacci G. 2019. Reference networks and civil codes. See Livermore & Rockmore 2019, pp. 339–65
- Bench-Capon T, ed. 1991. *Knowledge-Based Systems and Legal Applications*. San Diego, CA: Academic
- Bench-Capon T, Araszkievicz M, Ashley K, Atkinson K, Bex F, et al. 2012. A history of AI and law in 50 papers: 25 years of the International Conference on AI and Law. *Artif. Intell. Law* 20:215–319
- Betts KD, Jaep KD. 2017. The dawn of fully automated contract drafting: Machine learning breathes new life into a decades-old promise. *Duke Law Technol. Rev.* 15:216–33
- Beuve J, Moszoro MW, Saussier S. 2019. Political contestability and public contract rigidity: an analysis of procurement contracts. *J. Econ. Manag. Strategy* 28:316–35

⁴See, e.g., *Georgia v. Public.Resource.Org, Inc.*, 140 S. Ct. 1498 (2020) (holding that the official annotations of the State of Georgia cannot be copyrighted).

⁵Loi 2019–222 du 23 mars 2019 de programmation 2018–2022 et de réforme pour la justice (1) [Programming and Reform for Justice (1)], Article 33, Journal Officiel de la République Française [J.O.] [Official Gazette of France], March 24, 2019, texte no. 2, https://www.legifrance.gouv.fr/jo_pdf.do?id=JORFTEXT000038261631.

- Black RC, Hall MEK, Owens RJ, Ringsmuth EM. 2016. The role of emotional language in briefs before the US Supreme Court. *J. Law Courts* 4:377–407
- Black RC, Spriggs JF. 2008. An empirical analysis of the length of U.S. Supreme Court opinions. *Houst. Law Rev.* 45:621–82
- Black RC, Spriggs JF. 2013. The citation and deprecation of U.S. Supreme Court precedent. *J. Empir. Leg. Stud.* 10:325–58
- Black RC, Wedeking J. 2016. The influence of public sentiment on Supreme Court opinion clarity. *Law Soc. Rev.* 50:703–32
- Blei DM. 2012. Probabilistic topic models. *Commun. ACM* 55:77–84
- Blei DM, Lafferty J. 2007. A correlated topic model of science. *Ann. Appl. Stat.* 1:17–35
- Bommarito M, Katz DM. 2017. Measuring and modeling the U.S. regulatory ecosystem. *J. Stat. Phys.* 168:1125–35
- Bommarito MJ, Katz D, Zelner J. 2009. Law as a seamless web? Comparison of various network representations of the United States Supreme Court corpus (1791–2005). In *Proceedings of the 12th International Conference on Artificial Intelligence and Law*, pp. 234–35. New York: Assoc. Comput. Mach.
- Branting LK. 1991. Building explanations from rules and structured cases. *Int. J. Man-Mach. Stud.* 34:797–837
- Brüninghaus S, Ashley KD. 1997. Finding factors. Learning to classify case opinions under abstract fact categories. In *Proceedings of the 6th International Conference on Artificial Intelligence and Law*, pp. 123–31. New York: Assoc. Comput. Mach.
- Busch ML, Pelc KJ. 2019. Words matter: how WTO rulings handle controversy. *Int. Stud. Q.* 63:464–76
- Carlson K, Livermore MA, Rockmore D. 2016. A quantitative analysis of writing style on the U.S. Supreme Court. *Wash. Univ. Law Rev.* 93:1461–510
- Carlson K, Livermore MA, Rockmore D. 2020. The problem of data bias in the pool of published U.S. appellate court opinions. *J. Empir. Legal Stud.* 17:224–61
- Carmichael I, Wudel J, Kim M, Jushchuk J. 2017. Examining the evolution of legal precedent through citation network analysis. *N.C. Law Rev.* 96:227–69
- Carter DJ, Brown J, Rahmani A. 2016. Reading the High Court at a distance: topic modelling the legal subject matter and judicial activity of the High Court of Australia, 1903–2015. *Univ. N.S.W. Law J.* 39:1300–54
- Casey AJ, Niblett A. 2017. The death of rules and standards. *Indiana Law J.* 92:1401–47
- Caspi A, Stiglitz EH. 2020. Measuring discourse by algorithm. *Int. Rev. Law Econ.* 62:105863
- Choi SJ, Gulati GM. 2005. Which judges write their opinions (and should we care)? *Fla. State Univ. Law Rev.* 32:1076–122
- Choi SJ, Gulati GM. 2008a. Bias in judicial citations: A window into the behavior of judges? *J. Legal Stud.* 37:87–129
- Choi SJ, Gulati GM. 2008b. Trading votes for reasoning: covering in judicial opinions. *South. Calif. Law Rev.* 81:735–79
- Clark TS. 2017. Measuring law. In *Routledge Handbook of Judicial Behavior*, ed. RM Howard, KA Randazzo, pp. 84–94. New York: Routledge
- Clark TS, Lauderdale B. 2010. Locating Supreme Court opinions in doctrine space. *Am. J. Political Sci.* 54:871–90
- Clark TS, Lauderdale BE. 2012. The genealogy of law. *Political Anal.* 20:329–50
- Coglianese C. 2004. E-rulemaking: information technology and the regulatory process. *Adm. Law Rev.* 56:353–402
- Contissa G, Docter K, Lagioia F, Lippi M, Micklitz HW, et al. 2018. Automated processing of privacy policies under the EU General Data Protection Regulation. In *Legal Knowledge and Information Systems*, ed. M Palmirani, pp. 51–60. Amsterdam: IOS
- Contos G, Guyton J, Langetieg P, Vigil M. 2011. Individual taxpayer compliance burden: the role of assisted methods in taxpayers response to increasing complexity. In *IRS Research Bulletin: Proceedings of the IRS Research Conference 2010*, ed. ME Gangi, A Plumley, pp. 191–220. Washington, DC: Intern. Revenue Serv.
- Copus R, Hübert R, Laqueur H. 2019. Big data, machine learning, and the credibility revolution in empirical legal studies. See Livermore & Rockmore 2019, pp. 21–57

- Corley PC. 2008. The Supreme Court and opinion content: the influence of parties' briefs. *Political Res. Q.* 61:468–78
- Corley PC, Collins PM, Calvin B. 2011. Lower court influence on U.S. Supreme Court opinion content. *J. Politics* 73:31–44
- Cross FB. 2012. The ideology of Supreme Court opinions and citations. *Iowa Law Rev.* 97:693–751
- Cross FB, Pennebaker JW. 2014. The language of the Roberts Court. *Mich. State Law Rev.* 2014:853–94
- Cross FB, Spriggs JF. 2010. The most important (and best) Supreme Court opinions and justices. *Emory Law J.* 60:407–502
- Dadgostari F, Guim M, Beling PA, Livermore MA, Rockmore DN. 2020. Modeling law search as prediction. *Artif. Intell. Law.* <https://doi.org/10.1007/s10506-020-09261-5>
- Daniels JD, Rissland EL. 1997. Finding legally relevant passages in case opinions. In *Proceedings of the 6th International Conference on Artificial Intelligence and Law*, pp. 39–46. New York: Assoc. Comput. Mach.
- Derlén M, Lindholm J. 2013. Goodbye *van Gend en Loos*, hello *Bosman*? Using network analysis to measure the importance of individual CJEU judgments. *Eur. Law J.* 20:667–87
- Derlén M, Lindholm J. 2017. Is it good law? Network analysis and the CJEU's internal market jurisprudence. *J. Int. Econ. Law* 20:257–77
- Dumas M. 2019. Detecting ideology in judicial language. See Livermore & Rockmore 2019, pp. 383–405
- Evans M, McIntosh W, Lin J, Cates C. 2007. Recounting the courts? Applying automated content analysis to enhance empirical legal research. *J. Empir. Leg. Stud.* 4:1007–39
- Fagan F. 2015. From policy confusion to doctrinal clarity: successor liability from the perspective of big data. *Va. Law Bus. Rev.* 9:391–451
- Fagan F, Levmore S. 2019. The impact of artificial intelligence on rules, standards, and judicial discretion. *South. Calif. Law Rev.* 19:1–36
- Feldman A. 2017. A brief assessment of Supreme Court opinion language, 1946–2013. *Miss. Law J.* 86:105–49
- Feldman A. 2019. Opinion clarity in state and federal courts. See Livermore & Rockmore 2019, pp. 407–30
- Fowler JH, Jeon S. 2008. The authority of Supreme Court precedent. *Soc. Netw.* 30:16–30
- Fowler JH, Johnson TR, Spriggs JF, Jeon S, Wahlbeck PJ. 2007. Network analysis and the law: measuring the legal importance of precedents at the U.S. Supreme Court. *Political Anal.* 15:324–46
- Frankenreiter J. 2017. The politics of citations at the ECJ—policy preferences of EU member state governments and the citation behavior of judges at the European Court of Justice. *J. Empir. Legal Stud.* 14:813–57
- Frankenreiter J. 2019. Writing style and legal traditions. See Livermore & Rockmore 2019, pp. 153–90
- Funk K, Mullen LA. 2018. The spine of American law: digital text analysis and U.S. legal practice. *Am. Hist. Rev.* 123:132–64
- Gardner A. 1987. *An Artificial Intelligence Approach to Legal Reasoning*. Cambridge, MA: MIT Press
- Goelzhauser G, Cann DM. 2014. Judicial independence and opinion clarity on State Supreme Courts. *State Politics Policy Q.* 14:123–41
- Gonçalves T, Quaresma P. 2005. Is linguistic information relevant for the classification of legal texts? In *Proceedings of the Tenth International Conference on Artificial Intelligence and Law*, pp. 168–76. New York: Assoc. Comput. Mach.
- Grimmer J, Stewart BM. 2013. Text as data: the promise and pitfalls of automatic content analysis methods for political texts. *Political Anal.* 21:267–97
- Hadfield GK. 2008. Legal barriers to innovation: the growing economic cost of professional control over corporate legal markets. *Stanford Law Rev.* 60:1689–732
- Hausladen CI, Schubert MH, Ash E. 2020. Text classification of ideological direction in judicial opinions. *Int. Rev. Law Econ.* 62:105903
- Hinkle RK, Martin AD, Shaub JD, Tiller EH. 2012. A positive theory and empirical analysis of strategic word choice in district court opinions. *J. Legal Anal.* 4:407–44
- Hitt MP. 2016. Measuring precedent in a judicial hierarchy. *Law Soc. Rev.* 50:57–81
- Johnson SM. 2014. The changing discourse of the Supreme Court. *Univ. N.H. Law Rev.* 12:29–68
- Katz DM, Bommarito M. 2014. Measuring the complexity of the law: the United States Code. *Artif. Intell. Law* 22:337–74
- Katz DM, Bommarito MJ, Blackman J. 2017. A general approach for predicting the behavior of the Supreme Court of the United States. *PLOS ONE* 12:e0174698

- Kleinberg J, Ludwig J, Mullainathan S, Obermeyer Z. 2015. Prediction policy problems. *Am. Econ. Rev. Pap. Proc.* 105:491–95
- Klingenstein S, Hitchcock T, DeDeo S. 2014. The civilizing process in London's Old Bailey. *PNAS* 111:9419–24
- Kosnik L. 2014. Determinants of contract completeness: an environmental regulatory application. *Int. Rev. Law Econ.* 37:198–208
- Langford M, Behn D, Lie R. 2020. Stylometric analysis with machine learning: the case of investment treaty arbitration. In *Computational Legal Studies: The Promise and Challenge of Data-Driven Research*, ed. R Whalen. Cheltenham, UK: Edward Elgar. In press
- Larsson O, Naurin D, Derlén M, Lindholm J. 2017. Speaking law to power: the strategic use of precedent of the Court of Justice of the European Union. *Comp. Political Stud.* 50:879–907
- Lauderdale BE, Clark TS. 2014. Scaling politically meaningful dimensions using texts and votes. *Am. J. Political Sci.* 58:754–71
- Law DS. 2016. Constitutional archetypes. *Tex. Law Rev.* 95:153–243
- Law DS. 2018. The global language of human rights: a computational linguistic analysis. *Law Ethics Hum. Rights* 12:111–50
- Law DS, Zaring D. 2010. Law versus ideology: the Supreme Court and the use of legislative history. *William Mary Law Rev.* 51:1653–747
- Lee TR, Mouritsen SC. 2018. Judging ordinary meaning. *Yale Law J.* 127:788–879
- Leibon G, Livermore MA, Harder R, Riddell A, Rockmore D. 2018. Bending the law: geometric tools for quantifying influence in the multinet network of legal opinions. *Artif. Intell. Law* 26:145–67
- Leonard ME, Ross JV. 2016. Understanding the length of State Supreme Court opinions. *Am. Politics Res.* 44:710–33
- Li W, Azar P, Larochelle D, Hill P, Cox J, et al. 2013. Using algorithmic attribution techniques to determine authorship in unsigned judicial opinions. *Stanford Technol. Law Rev.* 16:503–33
- Lippi M, Palka P, Contissa G, Lagioia F, Micklitz HW, et al. 2017. Automated detection of unfair clauses in online consumer contracts. In *Legal Knowledge and Information Systems*, ed. A Wyner, G Casini, pp. 145–54. Amsterdam: IOS
- Livermore MA, Eidelman V, Grom B. 2018. Computationally assisted regulatory participation. *Notre Dame Law Rev.* 93:977–1034
- Livermore MA, Riddell AB, Rockmore DN. 2017. The Supreme Court and the judicial genre. *Ariz. Law Rev.* 59:837–901
- Livermore MA, Rockmore DL, eds. 2019. *Law as Data: Computation, Text, and the Future of Legal Analysis*. Santa Fe: Santa Fe Inst. Press
- Lupu Y, Fowler JH. 2013. Strategic citations to precedent on the U.S. Supreme Court. *J. Legal Stud.* 42:151–86
- Lupu Y, Voeten E. 2012. Precedent in international courts: a network analysis of case citations by the European Court of Human Rights. *Br. J. Political Sci.* 42:413–39
- Macey J, Mitts J. 2014. Finding order in the morass: the three real justifications for piercing the corporate veil. *Cornell Law Rev.* 100:99–156
- Marotta-Wurgler F, Taylor R. 2013. Set in stone? Change and innovation in consumer standard-form contracts. *NYU Law Rev.* 88:240–85
- McGuire KT, Vanberg G. 2005. *Mapping the policies of the U.S. Supreme Court: data, opinions, and constitutional law*. Paper presented at the Annual Meeting of the American Political Science Association, Washington, DC
- McLane J. 2019. Boilerplate and the impact of disclosure in securities dealmaking. *Vanderbilt Law Rev.* 72:191–295
- Mikolov T, Sutskever I, Chen K, Corrado GS, Dean J. 2013. Distributed representations of words and phrases and their compositionality. In *NIPS'13: Proceedings of the 26th International Conference on Neural Information Processing Systems*, Vol. 2, ed. CJC Burges, L Bottou, M Welling, Z Ghahramani, KQ Weinberger, pp. 3111–19. New York: Assoc. Comput. Mach.
- Miller JS. 2019. Law's semantic self-portrait: discerning doctrine with co-citation networks and keywords. *Univ. Pittsb. Law Rev.* 81:1–62

- Moszoro M, Spiller PT, Stolorz S. 2016. Rigidity of public contracts. *J. Empir. Legal Stud.* 13:396–427
- Moretti F. 2013. *Distant Reading*. London: Verso
- Mullainathan S, Spiess J. 2017. Machine learning: an applied econometric approach. *J. Econ. Perspect.* 31:87–106
- Niblett A. 2010. Do judges cherry pick precedents to justify extra-legal decisions? A statistical examination. *Md. Law Rev.* 70:101–38
- Niblett A, Yoon AH. 2015. Judicial disharmony: a study of dissent. *Int. Rev. Law Econ.* 42:60–71
- Nyarko J. 2019. We'll see you in... court! The lack of arbitration clauses in international contracts. *Int. Rev. Law Econ.* 58:6–24
- Nystrom EC, Tanenhaus DS. 2016a. The future of digital legal history: no magic, no silver bullets. *Am. J. Leg. Hist.* 56:150–67
- Nystrom EC, Tanenhaus DS. 2016b. "Let's change the law": Arkansas and the puzzle of juvenile justice reform in the 1990s. *Law Hist. Rev.* 34:957–97
- Oldfather CM, Bockhorst JP, Dimmer BP. 2012. Triangulating judicial responsiveness: automated content analysis, judicial opinions, and the methodology of legal scholarship. *Fla. Law Rev.* 64:1189–224
- Olsen HP, Küçüksu A. 2017. Finding hidden patterns in ECtHR's case law: on how citation network analysis can improve our knowledge of ECtHR's Article 14 practice. *Int. J. Discrim. Law* 17:4–22
- Owens RJ, Wedeking JP. 2011. Justices and legal clarity: analyzing the complexity of US Supreme Court opinions. *Law Soc. Rev.* 45:1027–61
- Patton D, Smith JL. 2017. Lawyer, interrupted: gender bias in oral arguments at the US Supreme Court. *J. Law Courts* 5:337–61
- Pelc KJ. 2014. The politics of precedent in international law: a social network application. *Am. Political Sci. Rev.* 108:546–64
- Petrov S. 2016. Announcing SyntaxNet: The world's most accurate parser goes open source. *Google Research Blog*, May 12. <https://research.googleblog.com/2016/05/announcing-syntaxnet-worlds-most.html>
- Potter RA. 2019. *Bending the Rules: Procedural Politicking in the Bureaucracy*. Chicago: Univ. Chicago Press
- Pozen DE, Talley EL, Nyarko J. 2019. A computational analysis of constitutional polarization. *Cornell Law Rev.* 105:1–84
- Quinn KM, Monroe BL, Colaresi M, Crespin MH, Radev DR. 2010. How to analyze political attention with minimal assumptions and costs. *Am. J. Political Sci.* 54:209–28
- Raskin M. 2017. The law and legality of smart contracts. *Georgetown Law Technol. Rev.* 2:305–28
- Rauterberg G, Talley E. 2017. Contracting out of the fiduciary duty of loyalty: an empirical analysis of corporate opportunity waivers. *Columbia Law Rev.* 117:1075–151
- Remus D, Levy F. 2017. Can robots be lawyers? Computers, lawyers, and the practice of law. *Georgetown J. Leg. Ethics* 30:501–58
- Rice D. 2014. The impact of Supreme Court activity on the judicial agenda. *Law Soc. Rev.* 48:63–90
- Rice D. 2017. Issue divisions and US Supreme Court decision making. *J. Politics* 79:210–22
- Rice D. 2019. Measuring the issue content of Supreme Court opinions. *J. Law Courts* 7:107–27
- Rice D, Rhodes JH, Nteta T. 2019. Racial bias in legal language. *Res. Politics* 6:1–7
- Rice D, Zorn C. 2019. Corpus-based dictionaries for sentiment analysis of specialized vocabularies. *Political Sci. Res. Meth.* 67:1–16
- Rissland EL, Ashley KD, Loui RP. 2003. AI and law: a fruitful synergy. *Artif. Intell.* 150:1–15
- Rissland EL, Skalak DB. 1991. CAARET: rule interpretation in a hybrid architecture. *Int. J. Man-Mach. Stud.* 34:839–87
- Roberts ME, Stewart BM, Nielsen RA. 2020. Adjusting for confounding with text matching. *Am. J. Political Sci.* In press
- Roberts ME, Stewart B, Tingley D. 2016. Navigating the local modes of big data: the case of topic models. In *Computational Social Science, Discovery and Prediction*, ed. RM Alvarez, pp. 51–97. New York: Cambridge Univ. Press
- Rockmore DN, Fang C, Foti NJ, Ginsburg T, Krakauer DC. 2017. The cultural evolution of national constitutions. *J. Assoc. Inf. Sci. Technol.* 69:483–94
- Romney CW. 2016. Using vector space models to understand the circulation of habeas corpus in Hawai'i. 1852–92. *Law Hist. Rev.* 34:999–1026

- Rosenthal JS, Yoon AH. 2011. Detecting multiple authorship of United States Supreme Court legal decisions using function words. *Ann. Appl. Stat.* 5:283–308
- Ruger TW, Kim PT, Martin AD, Quinn KM. 2004. The Supreme Court forecasting project: legal and political science approaches to predicting Supreme Court decisionmaking. *Columbia Law Rev.* 104:1150–209
- Ruhl JB, Katz DM. 2015. Measuring, monitoring, and managing legal complexity. *Iowa Law Rev.* 100:191–244
- Ruhl JB, Katz DM, Bommarito M. 2017. Harnessing legal complexity. *Science* 355:1377–78
- Ruhl JB, Nay J, Gilligan JM. 2018. Topic modeling the president: conventional and computational methods. *George Washington Law Rev.* 86:1243–315
- Sadl U, Olsen HP. 2017. Can quantitative methods complement doctrinal legal studies? Using citation network and corpus linguistic analysis to understand international courts. *Leiden J. Int. Law* 30:327–49
- Schmitz AJ. 2019. Expanding access to remedies through e-court initiatives. *Buffalo Law Rev.* 67:89–163
- Sergot MJ, Sadri F, Kowalsi RA, Kriwaczek F, Hammond P, Cory HT. 1986. The British Nationality Act as a logic program. *Comm. ACM* 29:370–86
- Smith JL. 2014. Law, fact, and the threat of reversal from above. *Am. Political Res.* 42:226–56
- Smith TA. 2007. The web of law. *San Diego Law Rev.* 44:309–54
- Stanford Nat. Lang. Process. Group. n.d. *The Stanford Parser: a statistical parser*. <https://nlp.stanford.edu/software/lex-parser.shtml>
- Stiglitz EH. 2014. Unaccountable midnight rulemaking? A normatively informative assessment. *Legis. Public Policy* 17:137–92
- Stiglitz EH. 2018. The limits of judicial control and the nondelegation doctrine. *J. Law Econ. Organ.* 34:27–53
- Strang LJ. 2016. How big data can increase originalism’s methodological rigor: using corpus linguistics to reveal original language conventions. *UC Davis Law Rev.* 50:1181–241
- Talley E, O’Kane D. 2012. The measure of a MAC: a machine-learning protocol for analyzing force majeure clauses in M&A agreements. *J. Inst. Theor. Econ.* 168:181–201
- Tarissan F, Nollez-Goldbach R. 2016. Analysing the first case of the International Criminal Court from a network-science perspective. *J. Complex Netw.* 4:616–34
- Thompson P. 2001. Automatic categorization of case law. In *Proceedings of the Eighth International Conference on Artificial Intelligence and Law*, ed. RP Loui, pp. 70–77. New York: Assoc. Comput. Mach.
- Tobia KP. 2020. Testing ordinary meaning: an experimental assessment of what dictionary definitions and linguistics usage data tell legal interpreters. *Harvard Law Rev.* 133. In press
- Varsava N. 2018. The elements of judicial style: a quantitative guide to Justice Gorsuch’s writing. *NYU Law Rev.* 93:75–112
- Wahlbeck PJ, Spriggs JF, Sigelman L. 2002. Ghostwriters on the court? A stylistic analysis of U.S. Supreme Court opinion drafts. *Am. Politics Res.* 30:166–92
- Whalen R. 2013. Modeling annual Supreme Court influence: the role of citation practices and judicial tenure in determining precedent network growth. In *Complex Networks*, ed. R Menezes, A Evsukoff, MC González, pp. 169–76. Berlin: Springer Verlag
- Whalen R. 2015. Judicial gobbledygook: the readability of Supreme Court writing. *Yale Law J. Forum* 125:200–11
- Whalen R, Uzzi B, Mukherjee S. 2017. Common Law evolution and judicial impact in the age of information. *Elon Law Rev.* 9:115–70
- Young DT. 2013. How do you measure a constitutional moment? Using algorithmic topic modeling to evaluate Bruce Ackerman’s theory of constitutional change. *Yale Law J.* 122:1990–2054