

Flexible, Model-Agnostic Method for Materials Data Extraction from Text Using General Purpose Language Models

Maciej P. Polak,^{*} Shrey Modi, Anna Latosinska, Jinming Zhang,
Ching-Wen Wang, Shanonan Wang, Ayan Deep Hazra, and Dane Morgan[†]
*Department of Materials Science and Engineering,
University of Wisconsin-Madison, Madison, Wisconsin 53706-1595, USA*
(Dated: February 13, 2023)

Accurate and comprehensive material databases extracted from research papers are critical for materials science and engineering but require significant human effort to develop. In this paper we present a simple method of extracting materials data from full texts of research papers suitable for quickly developing modest-sized databases. The method requires minimal to no coding, prior knowledge about the extracted property, or model training, and provides high recall and almost perfect precision in the resultant database. The method is fully automated except for one human-assisted step, which typically requires just a few hours of human labor. The method builds on top of natural language processing and large general language models but can work with almost any such model. The language models GPT-3/3.5, bart and DeBERTaV3 are evaluated here for comparison. We provide a detailed analysis of the methods performance in extracting bulk modulus data, obtaining up to 90% precision at 96% recall, depending on the amount of human effort involved. We then demonstrate the methods broader effectiveness by developing a database of critical cooling rates for metallic glasses.

I. INTRODUCTION

Obtaining reliable and comprehensive materials data is crucial for many research and industrial applications. If necessary information is not accessible through curated databases researchers typically must manually extract the data from research papers, a process that can be time-consuming and labor-intensive. Natural language processing (NLP) with general Language Models (LMs), and in particular, powerful large LMs (LLMs) trained of massive bodies of data, offer a new and potentially transformative technology for increasing the efficiency of extracting data from papers. These methods are particularly valuable when the data is embedded in the text of the documents, rather than being presented in a structured or tabular format, making it harder to find and extract. The rapid pace of development in NLP and frequent release of improved LLMs suggests they can be best utilized by methods which are easily adapted to new LLMs. In this paper we present such a flexible method for materials data extraction and demonstrate that it can achieve excellent precision and recall.

So far, the majority of materials data extraction approaches focus on fully automatic data extraction [1–3]. Automation is clearly desirable, particularly when extracting very large databases. However, more automation tends to require more complexity in the software, sophistication in training schemes, and knowledge about the extracted property. In addition, if a high level of completeness is required from a database, the recall of these approaches may not be sufficient. In such fully automated approaches a large amount of focus has been

placed on the complex task of named entity recognition (NER) [4–7], so that the property, material, values and units can be extracted accurately. However, automatic identification of an improper recognition is still very challenging, which can reduce the precision of such approaches. Tools for automatic materials and chemistry data extraction, like OSCAR4 [8] or ChemDataExtractor [9, 10], have been developed and used to successfully extract large databases. A recent example includes a database of over 22 thousand entries for relatively complex thermoelectric properties [11], at an average precision of 82.5% and a recall of 39.23%, or over 100 thousand band gap values [12], with an average precision of 84% and a recall of 64%. More complex information such as synthesis recipes [13–18] have also been extracted with automated NLP-based methods. Although not complete due to the relatively low recall, databases of that size are useful for training machine learning models [19–25], and would be very time consuming or impossible to extract with virtually any other method than full automation. Other recent examples of databases created in a similar way include photovoltaic properties and device material data for dye-sensitized solar cells [26], yield strength and grain size [27], and refractive index [28, 29]. Other notable databases gathered with NLP-based approaches include more complex information than just data values, such as synthesis procedures [14, 22]. Recently, another method for structured information extraction, making use of the GPT-3 capabilities was presented [30]. In that work, the focus is placed on the complicated NER tasks and relation extraction, at which GPT-3 excels. In that approach, more complex sentences can be successfully parsed into structured information. A "human in the loop" approach was used to fine-tune the model, a technique that seems to be emerging as a method of choice to obtain higher performing models. Impressive perfor-

^{*} mppolak@wisc.edu

[†] ddmorgan@wisc.edu

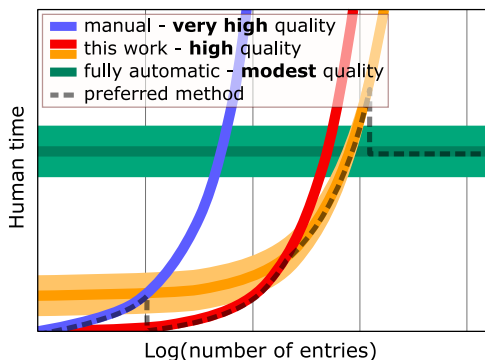


FIG. 1. Qualitative behavior of different types of approaches to data extraction, presented as human time required as a function of the size of the dataset. The broad range of the green (fully automatic), and orange (this work) represents the potential variation in the initial fixed time requirement, which may slightly influence the quality of the result. The dashed line suggests which method is the best choice for a given size of dataset.

mance was achieved in this work for structured information extraction, although at a price of a relatively large set of relatively complex training examples.

Depending on the nature of the data and the end goal for which the database is needed, there are different requirements for the resultant database and different optimal approaches for the data extraction. It is useful to organize methods along two broad axes. One axis is human time, which generally has the form

$$t = A + B \cdot N. \quad (1)$$

Here A is up-front fixed time to develop the method for a specific case and B is a marginal time cost whose contribution scales with some measure of data quantity N (N represents some function of the number of papers, sentences, and data to extract). The other axis is database quality, generally represented by some combination of precision (what fraction of the found data are right) and recall (what fraction of the available data was found) of the database. A schematic plot representing required human time as a function of the size of the database for different methods, and the quality of their results, is presented in Fig. 1.

One limiting case, which we will call "small data", is one where only a small amount of data is available in the literature (for example, properties that are very new, very hard to measure, or studied by only a small community), and where completeness and accuracy are highly valued. It is typical to develop a database for properties in the very small data limit fully manually, usually by experts in the field. Full manual curation is practical due to the limited number of papers and data and assures that the database is comprehensive and accurate. This fully manual approach is represented in Fig. 1 as the blue line, which is preferred (dashed) for very small datasets.

Even though it is technically slower than other methods, even in the very small data range, it is still the method of choice due to the highest possible quality of the results.

The opposite limiting case, which we will call "large data", is when there is a lot of data in the literature, the database is expected to be large, and modest precision and recall are acceptable. For example, such a database might be pursued for use for building machine learning regression and classification models on widely studied properties. For this large data case fully automated NLP-based approaches may be the most appropriate solution (see green curve in Fig. 1 for large number of entries). However, such an automated approach can result in an incomplete database that may not be sufficient for certain research or industrial applications, e.g., where extremes of performance of just a few materials might be the primary interest. In addition, fully automated approaches often require extensive retraining and building of parsers specific for different properties, as well as a significant amount of coding. These methods thus often require a significant initial investment of human time.

The logic of the best approaches for these extremes is simple. Large data (e.g., $> 10^4$ data points) can be most efficiently extracted by spending human time on automating the extraction (leading to large A and small B in Eq. 1), and reduced precision and recall is often of limited consequence since so much data is available. Small data sets (e.g., < 10 data points) can be most efficiently extracted by spending human time on directly extracting the data (leading to small A and large B in equation Eq. 1), and high precision and recall is typically more important for smaller databases. However, the optimal approach for the middle ground between these scales, which represents many databases in materials, is not obvious.

We propose a method that is most suited to creating these mid-size databases. Our method uses a combination of NLP methods, such as LLMs, with some degree of human supervision and input, which allows one to relatively quickly extract data of high quality while at the same time requiring minimal coding experience and up-front fixed human time cost. The method leads to modest A and B in Eq. 1, making it better than human extraction or full automation in the medium-data scale range. Two variants of the proposed methods are represented in Fig. 1 by the red and orange curves. They provide data of relatively high quality, approaching that of a fully manually created database, and scale well for medium sized databases. The proposed methods allow a database of up to the order of 1000 data points to be gathered in a few hours.

The approach taken in this work is to break up papers into sentences, use a LM to classify each sentence as relevant or not, parse each relevant category sentence with a LM into a structured set of target data, and then perform human review of the extracted structured data for validation and fixing. The LM classification is done either fully automatically (in a zero-shot fashion) or with some small human effort to fine-tune the LM with exam-

ple sentences. The LM classification step typically removes about 99% of the irrelevant data and leaves only about 1% to be further analyzed, dramatically reducing human labor. The final human review is very efficient as only highly structured data is presented, and most are already correct or nearly correct. This method results in an almost perfect precision and recall for the resultant database, comparable to a fully human curated database, but at 100 times or less human effort.

There are three major advantages of this method compared to other possible approaches of data extraction with NLP. First, the method is very easy to apply, requiring almost no coding or NLP expertise and very limited time from the user. For example, the case where the LM is provided by transformers zero-shot classification pipeline [31] requires just 3 lines of code that are provided on the huggingface website. As another example, in the case where the LLM GPT-3/3.5 is used, the API request is also just a few lines and provided to the user explicitly by the developers. Second, the method interfaces with the LM through a standard classification task available in any modern LM, making it possible to easily use the method with many present and likely any future LLMs. Thus the method can easily take advantage of the rapid improvements occurring in LLMs. Third, the method requires almost no knowledge about the property for which the data is to be extracted, with just the property name required for the basic application of the method.

In this paper we demonstrate our method by developing multiple databases with multiple LMs. The simplicity and flexibility of the method is illustrated by repeating the development of a benchmark bulk modulus database with multiple OpenAI GPT models, including the recently released GPT-3 and GPT-3.5 [32, 33], as well as the `bert`- and `DeBERTaV3`-based language models [34–37] hosted on huggingface, currently the most downloaded models for text classification. It is important to demonstrate the applicability of the method on both simple LMs that can be easily used on a personal computer, and on LLMs, which require significantly more computation and may be out of reach of most people’s resources for now. Even though there exist fully free and open LLMs, such as `OPT`[38] or `BLOOM`[39], their use is computationally expensive and not convenient, which contradicts the spirit of ease and accessibility of the presented method. Therefore, we opted for The OpenAI’s GPT-3/3.5 whose API allows one to efficiently use the LLM on outside servers, although is not free. GPT-3/3.5 is also currently the most popular LLM, so its a choice that will likely be relevant for many users.

We demonstrate and benchmark the method on raw texts of actual research papers, simulating how the method will likely be used by science and engineering communities. We first assess the precision and recall of the method on a small set of papers and the property bulk modulus in order to demonstrate the excellent accuracy that can be obtained with this method. We then

use the method to extract a modest sized but high quality database of critical cooling rates for metallic glasses.

The paper is organized as follows: Section II describes the methods in detail; Section III shows the results of benchmarks and statistical analysis of the obtained results for a Bulk Modulus database; Section IV discusses the database of critical cooling rates for metallic glasses, the possible utility of the method for purposes other than simple data extraction, and future possibilities in light of the rapid evolution of NLP methods and new LLMs. Section VI describes in detail the benchmark bulk modulus database used for assessment as well as the critical cooling rate of metallic glasses database.

II. DESCRIPTION OF THE METHOD

This section describes the steps involved in our method for data extraction from research papers. These steps are schematically presented in Fig. 2. Each step is first summarized (in bold), and then expanded to include details and observations we had during their development. Note that the first step is focused on gathering and basic processing of the papers and does not use the NLP methods. Furthermore, this step (or one very similar) is present in any data extraction method and is not specific to our method. Since this step is largely a universal preprocessing that does not have any direct connection to our method we label it the zeroth step, thereby allowing the first step to denote a step directly connected to our method.

STEP 0:

Obtaining and postprocessing the raw html/xml paper texts into human-readable format.

INPUT: This step starts with gathering the papers for analysis (e.g. from ScienceDirect API [40]) in an `xml/html` format. This usually involves searching for papers through a query to the publisher’s search engine, and simply downloading every paper that comes up as a result. If at this stage any relevant papers are missed, the data will not be extracted, so it is safer to use a broad search query or combine results of multiple queries. Including additional data does not increase the amount of human time, only involves more processing for the NLP model. This step takes very little human time, does not depend on the size of the database and is already mostly automated through the publisher, since only the query is required to obtain a list of matching documents.

SIMPLE TEXT PROCESSING: After the papers are downloaded the metadata and `html/xml` syntax is removed. We keep all the paragraphs and the title and remove the rest of the content. Then, we remove the `html/xml` markup syntax and tags. At this point all that is left is pure text. This cleaned up text is then split into separate sentences, according to regular rules

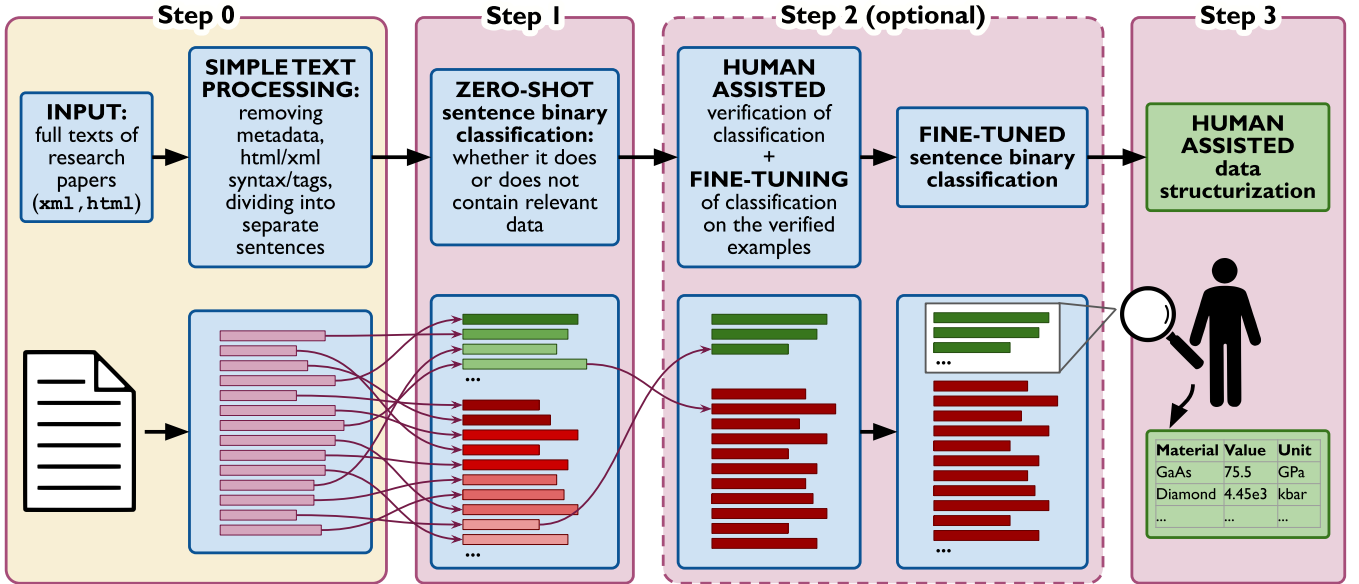


FIG. 2. A diagram of the steps necessary for NLP/LM data extraction in the proposed method. The process starts with gathering and preparing the documents to be analyzed, a process not involving any NLP (Step 0), then a LM is used to classify sentences by whether a sentence does or does not contain data for a given property (value and units) in a zero-shot fashion (Step 1). The pre-classified sentences are then (optionally) validated and used for fine-tuning the LM and reclassifying the sentences with higher quality (Step 2). Finally the data is structured by a LM/human assisted process, where the name of the material/system, the numerical value of the given property, its unit, and in some cases an additional detail, such as the temperature at which the value are obtained (Step 3). A detailed description of all steps can be found in Sec. II

for how sentences are terminated. At the end of this step we are left with the raw data that may be fed to the LM and analyzed. There is no need for any human evaluation of the data at this point. Whether an entire paper is unrelated, or some of the paragraphs contents are, it will simply be analyzed by the LM and deemed irrelevant. This step takes very little human time, and the amount required does not depend on the size of the database or the extracted property. The exact method for removal of `html/xml` syntax and splitting into sentences can vary. It can be done by text processing through regular expressions (an example can be found in the codes, Sec. VIII B), or ready-made specialized python libraries and their functions (such as `lxml`, `nltk.tokenize.sent_tokenize` [41, 42]), depending on the user preference.

It is important to note here that further simple text processing of the cleaned up text to keep only plausible sentences, e.g. using regular expressions to keep sentences with easily identified essential information, can, and probably should be performed at this point. Although such an additional processing step does not influence the method or the final outcome and quality of the produced database, such processing can significantly reduce the amount of data to be categorized by the LMs. This simple processing will certainly reduce the compute time needed for the NLP and can reduce costs if the LM is not free. How this text processing is performed depends on the task and the amount of knowledge about

the data or property to be extracted. For example, if we know the data is numeric we can keep just sentences containing a number. In the case of bulk modulus (see Dataset 1 in Sec. VI), keeping only sentences containing a number cuts the amount of data to be processed in half and does not lower recall (i.e., keeps all relevant sentences). If some amount of knowledge about the quantity to be extracted is available it can be used to further select the most promising sentences. For example, if we know the expected units of the data we can further process the remaining sentences to keep only those that contain such units. In the case of bulk modulus (see Dataset 1 in Sec. VI) keeping only sentences containing units of pressure (pascals and bars with possible metric prefixes, N/m^2), lowers the amount of possible candidates to less than 20% of the initial set, still without any loss of recall. Such refinements can be continued to further narrow down the search, but each subsequent step relies on a deeper knowledge of the property in question and increases the risk of reducing recall. In the work presented here we assume the most demanding situation for the method, in which no prior knowledge of the property is assumed. Therefore we only narrow down the search to sentences containing numerical values.

STEP 1:

Zero-Shot binary classification of sentences to produce unstructured data, i.e. set of sentences containing values for a given property. The classification puts sentences in two categories: *positive*, which are sentences containing the data (numerical value for a given property with its corresponding unit), and *negative*, which are sentences *not* containing the data.

Depending on the LLM used, the zero-shot may require as little as just the name of the desired property as the label of the class (name of the property), or require a full prompt phrase (e.g., GPT-3/3.5). Since the most recent and powerful LLMs make use of a prompt (e.g., GPT-3/3.5), we focus on that case here. The prompt (a single set of words, typically a phrase that makes grammatical sense) given to the model plays an important role. The impact of prompts has already been widely observed in NLP-based text to image generation tools (e.g., DALL-E2 [43], MidJourney [44], Stable Diffusion [45]) and a similar situation occurs in the present application. Depending on the completeness and phrasing of the prompt, the results for classification may be dramatically different. In our experience, however, prompts that do not contain false and misleading information almost always result high recall, and it is mainly the precision that is affected. In addition, more complex prompts do not necessarily guarantee a better result and may not be necessary.

Fig. 3 shows the zero-shot result statistics for the different models, including GPT-3/3.5 (whose technical names are `text-davinci-002` and `text-davinci-003`). The *p1* and *p2* stand for two different prompts.

p1: A sentence containing bulk modulus data must have its numerical value and the units of pressure. Does the following sentence contain bulk modulus data?

p2: Does the following sentence contain the value of bulk modulus?

Even though *p1* contains more seemingly valuable information, it did not necessarily perform better. We experimented with various different prompts, and straightforward prompts similar to *p2* performed most consistently and predictably. Therefore a simple prompt: *Does the following sentence contain the value of [name of property]?* is our strong recommendation.

It is worth noting that some models are not free to use, for example at the time of writing this article, the OpenAI models (GPT-3/3.5), after a small trial amount of free tokens, become paid on a per-token basis (1 token is approximately four English characters). Therefore, the flexibility to use different LMs presented by our method is very valuable, as some free models, while not necessarily capable of accurately performing the more complex tasks such as automated data structurization, and although overall generally less capable than GPT-3/3.5, perform well enough in the simple task of classification to produce satisfying results. However, in the case of

OpenAI GPT-3/3.5, both model usage and fine-tuning is done on outside servers, so in a situation where computational resources are not available to run locally, it may enable one to use the best models at a low cost.

STEP 2 (optional):

Human assisted verification of the zero-shot classification of sentences. The highest scoring unstructured data (most likely to be a true positive), pre-classified in Step 1, is manually classified into *positive* and *negative* sentences to provide a new training set of sentence for fine-tuning the classification process. The new training set is then used to fine-tune the model and classify the sentences again to obtain higher precision and recall.

This optional step is just a chance for the a human user to provide confirmation or correction to particularly important zero-shot classification data from Step 1 and then use those checks to fine-tune the LLM. Specifically, as the highest scoring sentences are being manually verified, a new training sets consisting of true positive and true negative examples is built. Since the precision of results of Step 1 is typically around 50% at 90% recall (see Fig. 3 (d)), the created sets are typically close to equal in size. The human labeled sets consist of *positive* cases, which represent *true positives* from Step 1, and *negative* cases, consisting of *false positives* from Step 1. The latter are the most valuable counter-examples for the negative training set, as these are the sentences the easiest for the model to confuse for positives. If after reaching the desired amount of verified positive sentences the corresponding set of negative sentences is smaller, it may be complemented with random sentences from the analyzed papers (the exceeding majority of which are negative). Fig. 4 shows how the classification model improves when fine-tuned on datasets of increasing size. A detailed analysis of that figure is present in Sec. III, where a conclusion is made that after around 100 positive sentences for the quicker learning models such as GPT-3/3.5 `davinci` or `bart`, we start to observe diminishing returns with this human labeled dataset size increase and it may not be worth spending more human time on obtaining more examples. Therefore, we recommend to perform the manual verification of the zero-shot classification until 100 positive sentences (and a corresponding 100 negative - made easy due to the 50% precision) are obtained, a number easy to remember and satisfactory for an efficient fine-tuning dataset.

This steps usually takes no more than 30 minutes - each sentence has to be classified only in a binary fashion which is a very simple task and takes just a few seconds per sentence. The classification is as simple and straightforward as assigning 1 for positive and 0 to each sentence in a spreadsheet. After this step is performed and the sentences are once again reclassified using the now fine-tuned model the precision and recall are greatly

improved, as can be seen in Fig. 4.

This Step 2 is optional and is generally done to improve the quality data collected from Step 1. Improving precision of data at this stage will reduce the human time needed in data structurization in Step 3 (see Sec. II) to review the data. However, for small data sets the human time in Step 3 is very modest, and this Step 2 may not be worth the extra effort. Thus whether it is performed or not typically depends on the size of the dataset. For small datasets, and if a recall of around 90% is satisfactory this step can be entirely omitted. As seen in Fig. 3 d), the precision at 90% recall after Step 1 is over 50% for the best models, which means that for every true positive sentence, there is only one false positive - a reasonable number to be removed by hand during data structurization (Step 3). For small datasets, up to a couple hundred values, verifying around 100 positive sentences to perform additional fine-tuning to improve the precision may turn out to be more labor intensive than proceeding straight to data structurization, and improving the precision manually by simply ignoring false positives. It is crucial to understand that the recall obtained at this step (or that has been obtained after Step 1, if this optional step is skipped) will be the recall of the final database, while the precision will be improved to near perfect in the next step (Step 3).

STEP 3:

Data structurization. In this step, extraction of the structured data is performed. Here, by structured data we mean the full information necessary to provide a datapoint: the name of the material/system, the numerical value of the given property, its unit, and in some cases an additional detail, such as the temperature at which the value was obtained. At the same time as the data is extracted the precision of the result is improved to perfect (or near perfect, depending on the expertise and accuracy of the human supervising this step), by simply ignoring the false positive sentences left over from previous steps. The result of this step is a final, curated structured database.

The user will typically perform this step by first ranking the sentences by their probability of being relevant, which is the output from Step 1 (or Step 2 if performed), and start reviewing the list at the top, working down until they decide to stop. As the user works through the results in that fashion, they traverse down the precision recall curve (PRC) (see. Figs. 3 (a) and 4 (a)). While the recall is impossible to assess without knowing the ground truth, the user is fully aware of the precision of the data they have already analyzed, therefore using the PRC they can estimate the recall and decide to stop when a desired recall is reached (with the assumption that the PRC are similar to those shown in Figs. 3 (a) and 4 (a)). For best models, reaching recall of around 90% (close to that of a

fully manual data curation) without performing the optional Step 2 happens for a precision close to 60%, while for a fine-tuned model (with Step 2), for a precision over 80%. The 90% recall seems to be a reasonable value to stop the process, as the precision sharply drops for higher values, which diminishes returns for the human time involved. However, this behavior may vary depending on the case, which will be discussed further in the Sec. III.

In general, human assisted data structurization, even when only the sentences containing the relevant data are given, may be a tedious and time consuming task. However, at this point it is the only method that can guarantee an almost perfect precision. For an unexperienced user, extracting one datapoint from a given sentence and its surrounding context fully by hand may take as long as 30 seconds, depending on the complexity of the property being analyzed and how it is typically expressed in research papers. Considering this, only relatively modest sized databases are reasonable to create. However, with experience, this time quickly reduces as the user gets used to the process. In addition, more experienced users may employ simple computer codes, e.g. based on regular expressions, which would preselect possible candidates for values and units, reducing the time significantly. In the longer term, it is likely the NLP tools will help automate this data structurization step, although they do not do this very effectively at present without either human supervision or a major effort to tune them. For example, GPT-3/3.5 is capable of parsing unstructured data in a zero-shot fashion, with no need for retraining. In the case of our Dataset 1 (bulk modulus) we found that in over 60% of cases GPT-3/3.5 is capable of correctly providing the entire data entry for a given property (name of the material/system, value, unit), and an incomplete datapoint (wrong material/system, but correct value and unit) in over 95% of cases. The only drawback that prohibits a full automation of this step with a LLM is the inability to automatically and unambiguously distinguish between a correct and incorrect extracted datapoints. However, human assistance in determining whether the data has been structurized properly, and in case it was not, fixing it by hand, can easily remedy that problem. Since almost all values and units are extracted correctly, and only less than half of the material names require fixing, using a LLM approach greatly reduces the human time and effort required for data structurization. Using an LLM we found that the average human time required to extract each good datapoint was reduced to under 10 seconds, keeping the same, almost perfect precision. Thus, and NLP-assisted data structurization, while still a tedious process, enables one create databases of up to around 1000 entries (more or less, depending on the users predisposition to and efficiency at repetitive tasks), in one workday. This timing includes the whole process, beginning (Step 0) to end (structurized database after Step 3), although almost all the human time is spent in Step 3.

While the value, units and the optional additional de-

tails most often occur within the positive sentence, the name of the material is often missing from that sentence (sentences are often similar to e.g. *We determined the bulk modulus to be 123 GPa.*). In those cases the system is described most often in the preceding sentence, and if not, then in the title of the paper. In a vast majority of cases (96% in our bulk modulus dataset) the full data information is available to be extracted from a sentence, that preceding it, and the title, so we do not search for it in other places. In the rare case when the full datapoint cannot be extracted, we record an incomplete datapoint. We also note that even in NLP models finely tuned for structured data extractions, the further apart the relevant data are from each other, the more difficult it is for the model to accurately extract the relevant data, so those datapoints would very likely be incomplete with other NLP-based approaches as well.

III. RESULTS

Fig. 3 summarizes the result from Step 1 in Sec. II for a bulk modulus analysis. The papers, sentences, ground truth category statistics, and other information is provided in Sec. VI. The ground truth for Step 1 was determined by human labeling. The following precision recall curves (PRCs) and receiver operating characteristic curves (ROCs) are constructed in the usual way, which is by plotting the relevant metrics while varying the cutoff used for the lowest value of probability accepted as a positive classification for sentence relevance. Panel (a) shows a PRC for the models tested in this paper. The two different curves for each of the **davinci** models correspond to two different prompts used in classification (see Sec. II). All of the tested models perform similarly, with **bart** struggling slightly more than others in achieving higher recall. A more quantitative measure of models' performance is presented in panel (b), where the area under the precision recall curve (AUC-PRC) alongside a maximum F1 score are presented. The **GPT-3/3.5 davinci** models, in particular using the first prompt (p1) show the highest scores, while **bart** and **DeBERTaV3** rank lowest in PRC-AUC. It is important to notice, however, that the datasets analyzed here are heavily imbalanced, with negative results outnumbering positives by more than 2 orders of magnitude. This places the naive *no skill* in Fig. 3 (a) line, representing an entirely random model, close to zero (as opposed to at 0.5 for a fully balanced set), lowering the entire PRC compared to a balance set. Fig. 3 (c) shows the ROC, which is insensitive to dataset imbalance, and shows much higher AUCs (panel (b)) than those of PRCs. The conclusions from ROCs are similar to those from PRCs; **GPT-3/3.5** performs best, with **bart** scoring lowest, while still performing reasonably well.

It is informative to consider the implication of the ROCs and PRCs for the efficiency of the human effort in our method. The step that requires most of the hu-

man time for a modest size database or larger (e.g., a few hundred entries or more) is Step 3, where the user must read and structure output from each sentence categorized as positive in Step 1 (or Step 2 if used). In Sec. II we suggested that the user limit their review of sentences in Step 3 unless a desired recall (implied by precision through the PRC) is achieved.

In some applications one might wish to target a high recall irrelevant of the human time required in step 3. To give a sense of how that might impact the method, Fig. 3 (d) and 4 (b) show the precision for 90% recall after Step 1 and after the optional Step 2, respectively. Consistent with our above discussions, the best models can achieve this recall with more than 50% precision using even just the zero-shot approach (Step 1). For less robust models, a 50% precision requires tuning (Step 2). For the worst models and using just zero-shot learning, the precision is about 17%, meaning the user would be extracting useful data from only about 1 in every 6 sentences reviewed. This would likely still be practical, but could become very tedious for a database of even a few hundred final entries. However, the important implication is that if one uses the best models (**GPT-3/3.5 davinci**), even a quite high recall requirement, e.g., 90%, can be achieved using very efficient sentence review, with almost every (more than 90%) sentence presented to the user containing relevant data.

Fig. 4 demonstrates how the performance of each of the models is improved if the optional fine-tuning in Step 2 is performed, as a function of the size of the training set. Panel (a) shows PRCs before fine-tuning (zero-shot) and compares them to PRCs after fine tuning on 100 and 200 positive sentences. While all models eventually show improvement, fine-tuning is clearly the most beneficial for the **GPT-3/3.5 davinci**. Similarly, various metrics describing the quality of the model are presented in Fig. 4 (b), where learning curves as a function of the size of the fine-tuning training set are shown. The x-axis represents the number of positive sentences included in the training set (with an assumed equal number of negative sentences). The shape of the learning curves differs for different models, with **GPT-3/3.5 davinci** model performing best (i.e. achieves higher performance metric values for smaller training sets) and learning the quickest (i.e. converges closer to best observed performance metric values for smaller training sets), **bart** following second, and **DeBERTaV3** third, across all metrics. For **davinci**, slope starts to decrease rapidly (curve starts to saturate) for as few as 60-80 positive sentences in the training set, for **bart** that occurs at around 100 positive sentences, and for **DeBERTaV3** closer to 160. Even though not all of the curves are fully saturated for the above mentioned dataset sizes, constructing larger fine-tuning training sets is likely to waste more human time than it is going to gain in Step 3. Our recommendation is to use a training set of around 100 positive sentences and the **GPT-3/3.5 davinci** model or the smaller and free **bart**. This size of 100 positive sentences is very

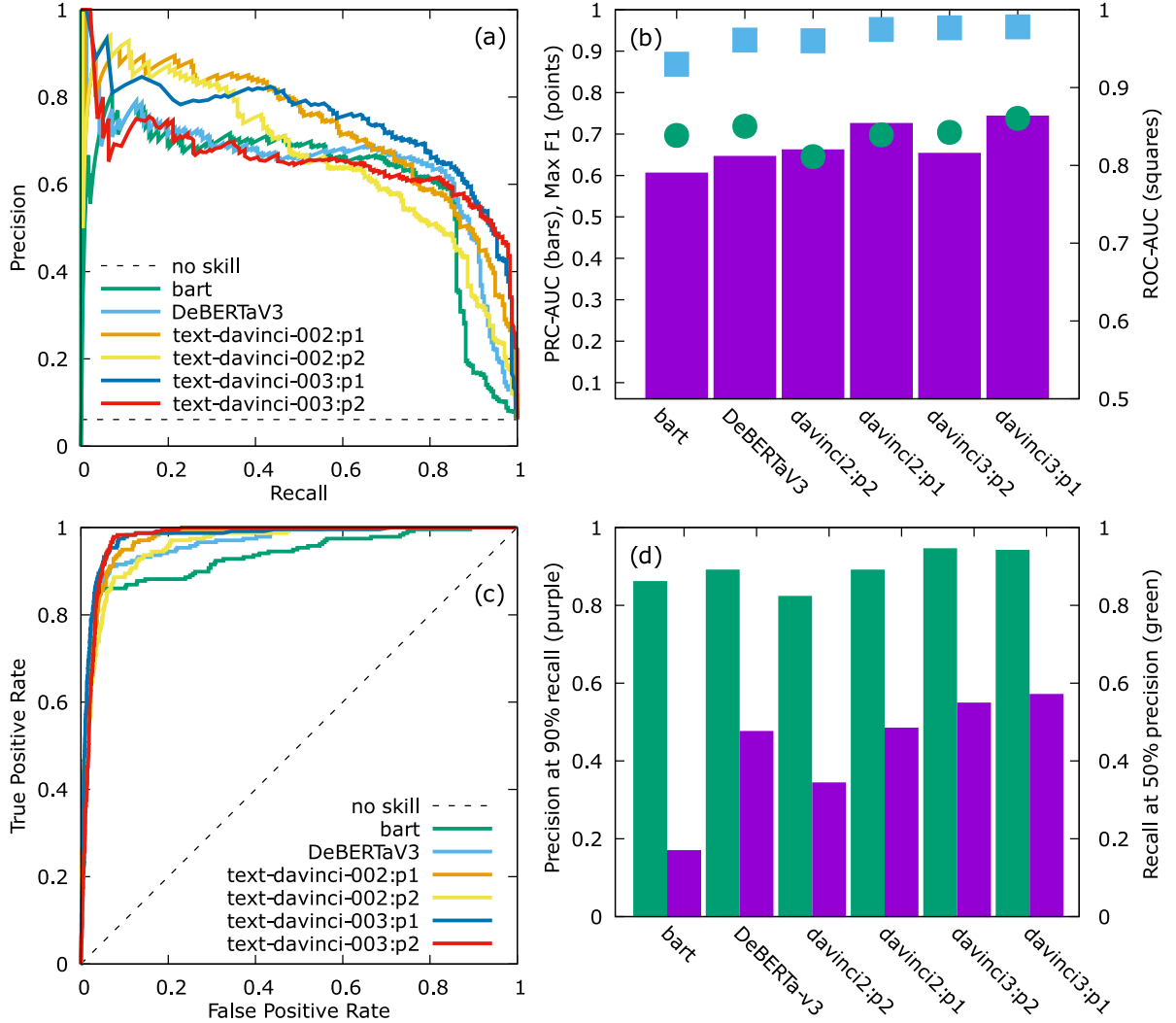


FIG. 3. Performance of different models after Step 1 (zero-shot binary classification of relevant sentences). (a) precision recall curves, (b) area under precision recall curve (PRC-AUC) (bars), maximum F1 score (circles), and area under receiver operating characteristic curves (ROC-AUC) (squares, right y-axis), (c) receiver operating characteristic curves, (d) precision at 90% recall and recall at 50% precision (right y-axis). The **no skill** line represent a baseline model where the classification is random. Labels in panels (b) and (d) have been shortened, but represent the same models as those in the legend of (a) and (c). **p1** and **p2** in the **davinci** models represent two different prompts (see Sec. II).

manageable to obtain with human-assisted verification of classification after Step 1, and typically does not take more than 30 minutes. Whether to perform the optional Step 2 (fine-tuning) will ultimately depend on the size of the database. As mentioned before, for larger databases, this improvement will be beneficial and save overall human time needed to curate the database by making Step 3 more efficient, while for small databases, up to a couple hundred datapoints, the time spent on the fine-tuning in Step 2 might be more than is saved during the data structurization in Step 3.

IV. DISCUSSION

To provide an example use-case for our method, we applied it to curate a high quality and highly accurate database of critical cooling rates for metallic glasses (Sec. VIB). 668 papers responded to the query "bulk metallic glass"+"critical cooling rate", which is more than what a human researcher would be analyze manually in a reasonable timeframe. The proposed method resulted in 443 datapoints consisting of the value of materials, their critical cooling rate, and the unit in which they were expressed in the paper. These results, include all mentions of critical cooling rates, with different degrees of specificity, e.g., accurate values for specific compositions (the ideal result), value ranges for specific materi-

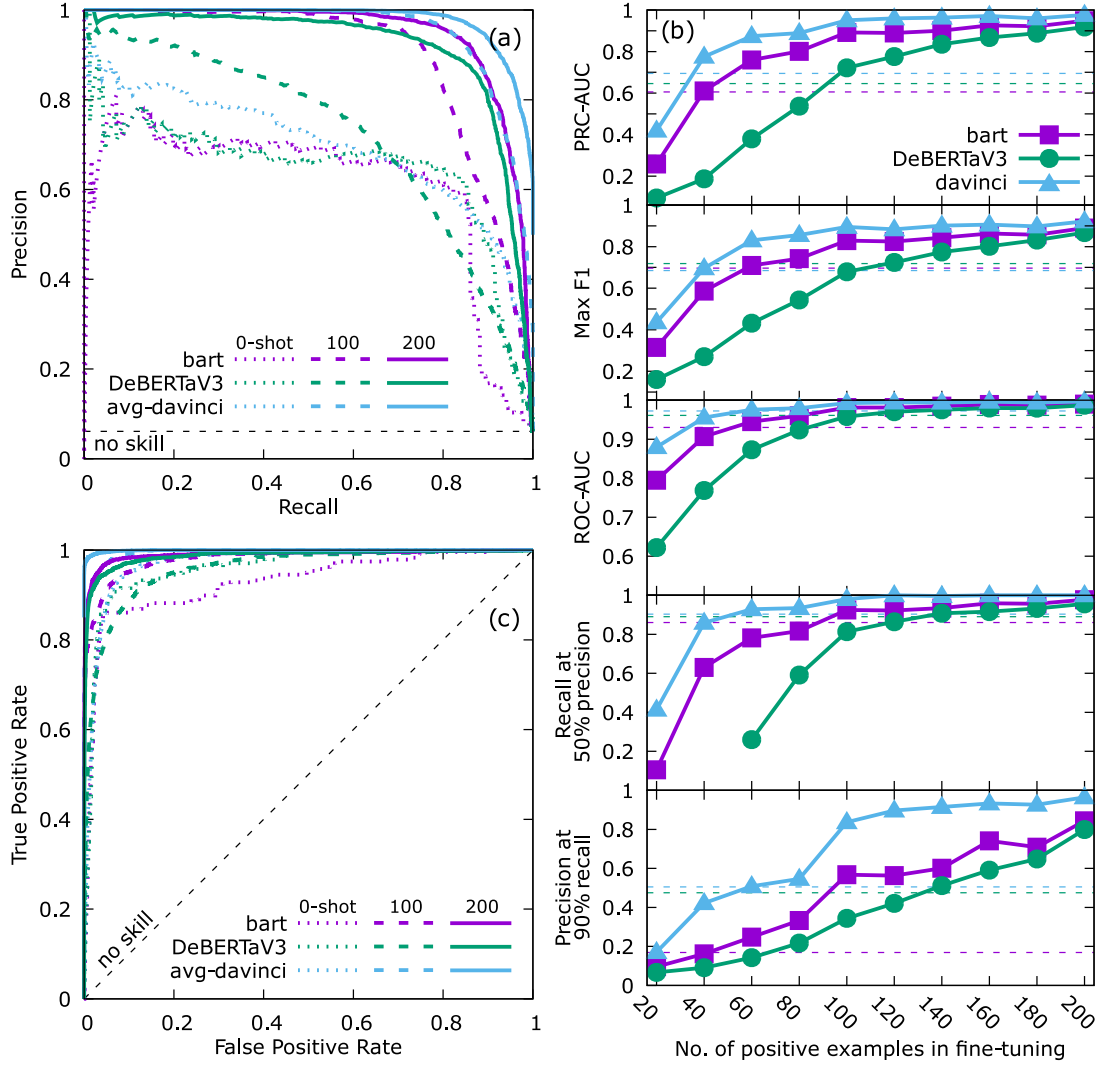


FIG. 4. Comparison of performance of different methods after fine-tuning (Step 2, Sec. II). Panel (a) shows precision recall curves, dotted lines correspond to the zero-shot (0-shot) result (davinci are averaged into one as described in the text), dashed and solid line correspond to fine-tuning on 100 and 200 positive sentence examples, respectively, (b) learning curves, i.e. performance metrics as a function of training set size, top to bottom: area under precision recall curve (PRC-AUC), maximum F1 score, area under receiver operating characteristic curves (ROC-AUC), recall at 50% precision, and precision at 90% recall. The horizontal thin dashed lines in corresponding colors represent zero-shot results. (c) receiver operating characteristic curves for the same data as in (a).

als, and value ranges for broad families of materials. The obtained database covers the range of expected values very well, with values ranging from 10^{-3} K/s for known bulk metallic glass formers, to 10^{11} K/s for particularly bad glass formers. The well known Pd-based bulk metallic glasses ($Pd_{43}Cu_2Ni_{10}P_{20}$ and $Pd_{43.2}Ni_{8.8}Cu_{28}P_{20}$) are identified as those with the lowest values of critical cooling rates, while simpler alloys such as AgCu, PdNi or NiBe and pure metals such as Co are identified as those with the highest critical cooling rates, which further validates the results. The obtained data, cleaned up for direct use in data oriented tasks (such as machine learning) i.e postprocessed to only include unique values for uniquely specified systems yielded 211 entries.

Within these, 129 are unique systems (multiple values are reported for some systems and we kept these to allow the user to manage them as they wish). The database is larger than the size of a recently published manually curated database of critical cooling rates [46], which is the most state-of-the-art and complete such database of which we are aware, and consists of only 77 unique compound datapoints.

Searching for a given property does not typically add any restrictions on the search other than the property itself, i.e., the search is unrestricted. In the case of the method proposed here, unrestricted search will identify and help extract all datapoints for the target property from the input set of documents. Therefore, if the user

desires a database limited to, for example, a given family of systems, the limitation would have to be enforced in some additional way. This constraint could be done by limiting the input set of documents through a more strict search query, but even that does not guarantee that only the desired values will be extracted, as many papers mention a wide range of results, even if technically focused on a particular topic. Limiting the final database can be easily done manually in Step 3 (Sec. II), but depending on the property and the size of the desired subset, limiting the data at that stage may take a lot of human time and be inefficient. In principle, more restrictions than just the property can be imposed on the NLP level, but such abstract concepts as families of materials are very challenging even for the best LLMs and greatly reduce the quality of the zero-shot results (Sec. II) and would require significantly more training in (Sec. II). This problem is highly dependent on the property in question. For example, an unrestricted search for critical cooling rates while limiting the search in Step 0 to papers responding to a query "bulk metallic glasses"+"critical cooling rate" was quite effective for our goals of obtaining all ranges of critical cooling rates for metallic glasses. But if one wanted, say, an overpotential for water splitting, restrictions on many factors, e.g., temperature or pH, might be essential to obtaining useful result and difficult to screen on in the initial Step 0.

A particular example of where unrestricted searches can be problematic occurs when searching for properties which are relevant in many fields when one is interested in only a particular field, and/or which have many possible associated restrictions which are needed to make the data useful. A specific example of this problem occurred for us when we explored constructing a database for "area specific resistance" (ASR) for anode materials of proton conducting cells. In step zero we searched for "area specific resistance"+"proton conducting fuel cells (and similar terms)" The method proved very successful at identifying sentences containing ASR and structuring the data, as it was asked to do. However, the method captured ASR in a wide variety of contexts, including single phase and composite materials, porous and non-porous materials, electrodes and electrolytes, steels, interconnects, coatings, varying temperatures, and ASR in both fuel cell and electrolysis operation modes. To obtain a simple and immediately useful dataset we were interested in single phase dense anodes operating in fuel cell mode with temperature information. Imposing such limitations was dramatically harder than the basic data extraction. Although one might have different goals than the ones just mentioned, it is very unlikely that one is interested in gathering information for all of the above data in a single database. Restricting the set of input documents was able to help to a certain degree to move the balance of the obtained results in the desired direction, but did not solve the issue entirely. From such a wide variety of contexts, identifying only those we were interested in required a relatively deep knowledge from

the person performing the data extraction and required significantly more human time to extract than in case of datasets where the property is more uniquely identified. In fact, we stopped developing this database due to these many challenges, although for someone willing to commit 4-5 days of human time in step 3 the desired database is certainly practical to develop.

On the other hand, the lack of restrictions in the model may have other benefits, as it expands the possibilities of the kinds of information that can be extracted. For example, our method can be used to extract many kinds of information, not just property values. Step 1 with models like GPT-3/3.5 broadly describe the type of text we are looking for, and Step 2 fine-tunes to better classify the relevant sentences. While we utilized this classification search to find sentences containing numerical data for a given property, it can be used to search for more abstract concepts, such as suitability of a given material for a certain application, personal opinions of authors about promising directions of future research, or any other concept that can be characterized as a group of example texts for the model to train on, and classify in a binary fashion.

It is also important to remember that the method we present here is not restricted to the LMs explored in this paper, and is in fact designed to be quickly adapted to new and improved LMs.

V. SUMMARY

We have shown a simple and efficient method for materials data extraction from research papers. The simple NLP concept of binary text classification is involved as a key step in the method, which allows for a high flexibility in the language model used as virtually all modern language models are very capable at text classification. We determined GPT-3/3.5 davinci to be the best performer, but evaluated other alternatives such as bart or DeBERTaV3. By including a highly-optimized human-assisted step in the process, we minimized the amount of coding and prior knowledge about the extracted property necessary to achieve a high recall and nearly perfect precision. A modest sized database of up to around 1000 entries can be extracted in around one workday with this method. The method is assessed vs. ground truth on a bulk modulus database and then applied to construct curated database of critical cooling rate of metallic glasses.

VI. DATASETS

Below, the details about the datasets are provided. As a result of this paper a high quality database of critical cooling rates for bulk metallic glasses has been curated, as well as a benchmark-only dataset - the bulk modulus dataset, which was used to assess the model. Information on accessing the datasets can be found in Sec. VIII B.

A. Bulk modulus

The bulk modulus is a benchmark dataset is a test dataset consisting of 100 papers randomly chosen from the over 10000 paper results of a search query "*bulk modulus*" + "*crystal*" returned from the ScienceDirect API. In the written text of these 100 papers, there are 18408 sentences in total, out of which 237 sentences mention the value of bulk modulus. This dataset is used as a benchmark so a human ground truth is extracted. To avoid excessive time spent establishing the human labeled ground truth this database uses only a very small fraction of the total available papers. Thus the bulk modulus database is not nearly complete.

For the zero-shot case (only step 1 and not step 2) the approach effectively has no training data and can just be assessed on the test data described above. However, when step 2 included the fine-tuning requires additional data (effectively a training data set). For this fine-tuning process an additional 339 positive and 484 negative sentences have been extracted from papers not included in the 100 papers in the test set. These additional sentences are use to investigate how fine-tuning improves the model and plot learning curves (see Fig. 4).

B. Critical cooling rates for bulk metallic glasses

This dataset consists of data gathered from 668 papers based on a result of a search query "bulk metallic glass" + "critical cooling rate" from Elsevier's ScienceDirect API. These papers consisted of 107386 sentences, out of which 347 were identified as positive (containing values of critical cooling rates), after applying the workflow described in Sec. II, including the optional Step 2 in order to provide best quality data. From these 347 sentences, 443 critical cooling rate data points (consisting of the material name, critical cooling rate value and units) were extracted and are collected as a final database presented. Additionally, that data was manually postprocessed to include only unique datapoints (removing duplicate results, i.e. the same values reported in multiple papers), remove those which included ranges or limits or values, or where the material's composition was not explicitly given, and unify the formatting of the materials compositions, which resulted in 211 unique datapoints.

ACKNOWLEDGMENTS

The research in this work was primarily supported by the National Science Foundation Cyberinfrastructure for Sustained Scientific Innovation (CSSI) Award No. 1931298. Additional support for engaging undergraduates in research was provided by the National Science Foundation Training-based Workforce Development for Advanced Cyberinfrastructure (CyberTraining)

Award No. 2017072. We thank Harmon Bhasin and Shanchao Liang for their participation in this project.

VII. CONTRIBUTIONS

M. P. P. conceived the study, performed the modeling, tests and prepared/analyzed the results, S. M., J. Z., J. W., S. W. and A. D. H. performed tests and assessed methods as a part of the Informatics Skunkworks program for undergraduate researchers, A. L. post-processed and extracted data, D. M. guided and supervised the research. Writing of the manuscript was done by M. P. P. and D. M.. All authors read, revised and approved the final manuscript.

VIII. METHODS

As with any machine learning model, there are hyperparameters that may be optimized. Our experience showed that there is very little to be gained by performing the costly optimization, and throughout the paper we used default recommended values for all models. In OpenAI GPT-3/3.5 *davinci* we used the default values for fine-tuning, and when using both the pretrained *text-davinci-002/003* and fine-tuned *davinci* we set the frequency and presence penalties as well as temperature to 0. The fine-tuning of *bart* and *DeBERTaV3* was performed with default recommended values too, which is a learning rate of $2e - 5$, batch size of 16, 5 epochs, and 0.01 weight decay. Full and detailed input files can be found in [47]. Python codes were executed with *Python* ver. 3.10.6.

A. Definition of statistical quantities

True positive (TP) - a sentence containing numerical data for a given property.

True negative (TN) - a sentence not containing numerical data for a given property.

False positive (FP) - a sentence not containing the numerical data for a given property but is identified as one that does.

False negative (FN) - a sentence containing numerical data for a given property but is identified as one that does not.

Precision:

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Recall (True Positive Rate):

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

False Positive Rate (FPR):

$$FPR = \frac{FP}{FP + TN} \quad (4)$$

F1 score:

$$F1 = \frac{2TP}{2TP + FP + FN} \quad (5)$$

B. Data Availability

The databases curated as a result of this paper, all datasets used in the assessment of the

method, as well as the codes and software used in this paper are available on figshare [47]: <https://doi.org/10.6084/m9.figshare.21861948>. The codes are included for full transparency, but were developed for internal use only, so they contain very limited error handling, and the authors do not guarantee that they will work universally on every system. All parameters used for the model fine-tuning and zero-shot classification can be found in the codes.

C. Competing Interest

The authors declare no competing interest.

-
- [1] E. A. Olivetti, J. M. Cole, E. Kim, O. Kononova, G. Ceder, T. Y.-J. Han, and A. M. Hiszpanski, Data-driven materials research enabled by natural language processing and information extraction, *Applied Physics Reviews* **7**, 041317 (2020).
 - [2] O. Kononova, T. He, H. Huo, A. Trewartha, E. A. Olivetti, and G. Ceder, Opportunities and challenges of text mining in materials research, *iScience* **24**, 102155 (2021).
 - [3] M. Krallinger, O. Rabal, A. Lourenço, J. Oyarzabal, and A. Valencia, Information retrieval and text mining technologies for chemistry, *Chemical Reviews* **117**, 7673 (2017).
 - [4] A. Trewartha, N. Walker, H. Huo, S. Lee, K. Cruse, J. Dagdelen, A. Dunn, K. A. Persson, G. Ceder, and A. Jain, Quantifying the advantage of domain-specific pre-training on named entity recognition tasks in materials science, *Patterns* **3**, 100488 (2022).
 - [5] L. Weston, V. Tshitoyan, J. Dagdelen, O. Kononova, A. Trewartha, K. A. Persson, G. Ceder, and A. Jain, Named entity recognition and normalization applied to large-scale information extraction from the materials science literature, *Journal of Chemical Information and Modeling* **59**, 3692 (2019).
 - [6] X. Zhao, J. Greenberg, Y. An, and X. T. Hu, Fine-tuning bert model for materials named entity recognition, *2021 IEEE International Conference on Big Data (Big Data)*, 3717 (2021).
 - [7] T. Isazawa and J. M. Cole, Single model for organic and inorganic chemical named entity recognition in chemdataextractor, *Journal of Chemical Information and Modeling* **62**, 1207 (2022).
 - [8] D. M. Jessop, S. E. Adams, E. L. Willighagen, L. Hawizy, and P. Murray-Rust, OSCAR4: a flexible architecture for chemical text-mining, *Journal of Cheminformatics* **3**, 41 (2011).
 - [9] M. C. Swain and J. M. Cole, Chemdataextractor: A toolkit for automated extraction of chemical information from the scientific literature, *Journal of Chemical Information and Modeling* **56**, 1894 (2016).
 - [10] J. Mavračić, C. J. Court, T. Isazawa, S. R. Elliott, and J. M. Cole, Chemdataextractor 2.0: Autopopulated ontologies for materials science, *Journal of Chemical Information and Modeling* **61**, 4280 (2021).
 - [11] O. Sierpeklis and J. Cole, A thermoelectric materials database auto-generated from the scientific literature using chemdataextractor, *Sci Data* **9**, 648 (2022).
 - [12] Q. Dong and J. Cole, Auto-generated database of semiconductor band gaps using chemdataextractor, *Sci Data* **9**, 193 (2022).
 - [13] O. Kononova, H. Huo, T. He, Z. Rong, T. Botari, W. Sun, V. Tshitoyan, and G. Ceder, Text-mined dataset of inorganic materials synthesis recipes, *Scientific Data* **6**, 203 (2019).
 - [14] Z. Wang, O. Kononova, K. Cruse, *et al.*, Dataset of solution-based inorganic materials synthesis procedures extracted from the scientific literature, *Sci Data* **9**, 231 (2022).
 - [15] Z. Wang, O. Kononova, K. Cruse, T. He, H. Huo, Y. Fei, Y. Zeng, Y. Sun, Z. Cai, W. Sun, and G. Ceder, Dataset of solution-based inorganic materials synthesis procedures extracted from the scientific literature, *Scientific Data* **9**, 231 (2022).
 - [16] E. Kim, K. Huang, A. Saunders, A. McCallum, G. Ceder, and E. Olivetti, Materials synthesis insights from scientific literature via text extraction and machine learning, *Chemistry of Materials* **29**, 9436 (2017).
 - [17] Z. Jensen, E. Kim, S. Kwon, T. Z. H. Gani, Y. Román-Leshkov, M. Moliner, A. Corma, and E. Olivetti, A machine learning approach to zeolite synthesis enabled by automatic literature data extraction, *ACS Central Science* **5**, 892 (2019).
 - [18] E. Kim, Z. Jensen, A. van Grootel, K. Huang, M. Staib, S. Mysore, H.-S. Chang, E. Strubell, A. McCallum, S. Jegelka, and E. Olivetti, Inorganic materials synthesis planning with literature-trained neural networks, *Journal of Chemical Information and Modeling* **60**, 1194 (2020).
 - [19] D. Morgan and R. Jacobs, Opportunities and challenges for machine learning in materials science, *Annual Review of Materials Research* **50**, 71 (2020).
 - [20] J. E. Saal, A. O. Oliynyk, and B. Meredig, Machine learning in materials discovery: Confirmed predictions and their underlying approaches, *Annual Review of Materials Research* **50**, 49 (2020).
 - [21] C. Court and J. Cole, Magnetic and superconducting phase diagrams and transition temperatures predicted using text mining and machine learning, *npj Comput Mater* **6**, 18 (2020).

- [22] H. Huo, C. J. Bartel, T. He, A. Trewartha, A. Dunn, B. Ouyang, A. Jain, and G. Ceder, Machine-learning rationalization and prediction of solid-state synthesis conditions, *Chemistry of Materials* **34**, 7323 (2022).
- [23] J. Zhao and J. M. Cole, Reconstructing chromatic-dispersion relations and predicting refractive indices using text mining and machine learning, *Journal of Chemical Information and Modeling* **62**, 2670 (2022).
- [24] C. Karpovich, Z. Jensen, V. Venugopal, and E. Olivetti, Inorganic synthesis reaction condition prediction with generative machine learning 10.48550/ARXIV.2112.09612 (2021).
- [25] A. B. Georgescu, P. Ren, A. R. Toland, S. Zhang, K. D. Miller, D. W. Apley, E. A. Olivetti, N. Wagner, and J. M. Rondinelli, Database, features, and machine learning model to identify thermally driven metal-insulator transition compounds, *Chemistry of Materials* **33**, 5591 (2021).
- [26] E. Beard and J. Cole, Perovskite- and dye-sensitized solar-cell device databases auto-generated using chemdataextractor, *Sci Data* **9**, 329 (2022).
- [27] P. Kumar, S. Kabra, and J. Cole, Auto-generating databases of yield strength and grain size using chemdataextractor, *Sci Data* **9**, 292 (2022).
- [28] J. Zhao and J. M. Cole, Reconstructing chromatic-dispersion relations and predicting refractive indices using text mining and machine learning, *Journal of Chemical Information and Modeling* **62**, 2670 (2022).
- [29] J. Zhao and J. Cole, A database of refractive indices and dielectric constants auto-generated using chemdataextractor, *Sci Data* **9**, 192 (2022).
- [30] A. Dunn, J. Dagdelen, N. Walker, S. Lee, A. S. Rosen, G. Ceder, K. Persson, and A. Jain, Structured information extraction from complex scientific text with fine-tuned large language models 10.48550/ARXIV.2212.05238 (2022).
- [31] T. Wolf, L. Debut, V. Sanh, J. Chaumond, C. Delangue, A. Moi, P. Cistac, T. Rault, R. Louf, M. Funtowicz, J. Davison, S. Shleifer, P. von Platen, C. Ma, Y. Jernite, J. Plu, C. Xu, T. L. Scao, S. Gugger, M. Drame, Q. Lhoest, and A. M. Rush, Transformers: State-of-the-art natural language processing, in *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations* (Association for Computational Linguistics, Online, 2020) pp. 38–45.
- [32] T. B. Brown, B. Mann, N. Ryder, M. Subbiah, J. Kaplan, P. Dhariwal, A. Neelakantan, P. Shyam, G. Sastry, A. Askell, S. Agarwal, A. Herbert-Voss, G. Krueger, T. Henighan, R. Child, A. Ramesh, D. M. Ziegler, J. Wu, C. Winter, C. Hesse, M. Chen, E. Sigler, M. Litwin, S. Gray, B. Chess, J. Clark, C. Berner, S. McCandlish, A. Radford, I. Sutskever, and D. Amodei, Language models are few-shot learners 10.48550/ARXIV.2005.14165 (2020).
- [33] L. Ouyang, J. Wu, X. Jiang, D. Almeida, C. L. Wainwright, P. Mishkin, C. Zhang, S. Agarwal, K. Slama, A. Ray, J. Schulman, J. Hilton, F. Kelton, L. Miller, M. Simens, A. Askell, P. Welinder, P. Christiano, J. Leike, and R. Lowe, Training language models to follow instructions with human feedback 10.48550/ARXIV.2203.02155 (2022).
- [34] M. Lewis, Y. Liu, N. Goyal, M. Ghazvininejad, A. Mohamed, O. Levy, V. Stoyanov, and L. Zettlemoyer, Bart: Denoising sequence-to-sequence pre-training for natural language generation, translation, and comprehension 10.48550/ARXIV.1910.13461 (2019).
- [35] W. Yin, J. Hay, and D. Roth, Benchmarking zero-shot text classification: Datasets, evaluation and entailment approach 10.48550/ARXIV.1909.00161 (2019).
- [36] P. He, J. Gao, and W. Chen, Debertav3: Improving deberta using electra-style pre-training with gradient-disentangled embedding sharing 10.48550/ARXIV.2111.09543 (2021).
- [37] A. Conneau, G. Lample, R. Rinott, A. Williams, S. R. Bowman, H. Schwenk, and V. Stoyanov, Xnli: Evaluating cross-lingual sentence representations 10.48550/ARXIV.1809.05053 (2018).
- [38] S. Zhang, S. Roller, N. Goyal, M. Artetxe, M. Chen, S. Chen, C. Dewan, M. Diab, X. Li, X. V. Lin, T. Mihaylov, M. Ott, S. Shleifer, K. Shuster, D. Simig, P. S. Koura, A. Sridhar, T. Wang, and L. Zettlemoyer, Opt: Open pre-trained transformer language models 10.48550/ARXIV.2205.01068 (2022).
- [39] B. Workshop, :, T. L. Scao, A. Fan, C. Akiki, E. Pavlick, S. Ilić, D. Hesslow, R. Castagné, A. S. Luccioni, F. Yvon, and M. Gallé et al., Bloom: A 176b-parameter open-access multilingual language model 10.48550/ARXIV.2211.05100 (2022).
- [40] Elsevier developer portal, <https://dev.elsevier.com>, [Online; accessed 08-Feb-2023].
- [41] S. Bird, E. Klein, and E. Loper, *Natural Language Processing with Python: Analyzing Text with the Natural Language Toolkit* (O'Reilly, 2009).
- [42] S. Behnel, M. Faassen, and I. Bicking, lxml: Xml and html with python (2005).
- [43] A. Ramesh, P. Dhariwal, A. Nichol, C. Chu, and M. Chen, Hierarchical text-conditional image generation with clip latents 10.48550/ARXIV.2204.06125 (2022).
- [44] Midjourney, <https://www.midjourney.com>, [Online; accessed 08-Feb-2023].
- [45] R. Rombach, A. Blattmann, D. Lorenz, P. Esser, and B. Ommer, High-resolution image synthesis with latent diffusion models, in *2022 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)* (2022) pp. 10674–10685.
- [46] B. T. Afflerbach, C. Francis, L. E. Schultz, J. Spethson, V. Meschke, E. Strand, L. Ward, J. H. Perepezko, D. Thoma, P. M. Voyles, I. Szlufarska, and D. Morgan, Machine learning prediction of the critical cooling rate for metallic glasses from expanded datasets and elemental features, *Chemistry of Materials* **34**, 2945 (2022).
- [47] M. P. Polak, S. Modi, A. Latosinska, J. Zhang, J. Wang, S. Wang, A. D. Hazra, and D. Morgan, Data for "Flexible, Model-agnostic Method for Data Extraction from Text Using General Purpose Natural Language Processing" 10.6084/m9.figshare.21861948 (2023).