

Occupational Classifications: A Machine Learning Approach

Akina Ikudo, University of California, Los Angeles

Julia Lane, New York University and U.S. Census Bureau

Joseph Staudt, U.S. Census Bureau

Bruce Weinberg, Ohio State University

Characterizing people's occupations is important for both policy and research. Because administrative records, rather than survey data, are increasingly being used to describe labor market activity, it will become important to find new low-cost approaches to describing occupations. We apply new techniques—machine learning—to new sources of data to investigate the potential of using algorithms to classify occupations. Our analysis of human resource data from large firms (universities) and the US Census Bureau shows that occupational classifications are inherently noisy; job titles by themselves have insufficient consistency across institutions to serve as the sole basis for the reliable assignment of occupations. We do find, however, that a large number of relatively sparsely populated job titles can be assigned algorithmically, which could greatly reduce cost with little impact on accuracy.

I. Introduction

Characterizing the work that people do on their jobs is a long-standing and core issue in survey research. Traditionally, classification has been done manually, but there is an extensive body of literature on the associated challenges, well summarized in an influential paper by Mellow and Sider (1) and in a later paper by Mathiowetz (2). Many survey organizations are beginning to investigate the potential of using new computational tools to automatically classify workers' occupations.

At the same time, there has been a surge of interest in using administrative wage records to directly capture occupations in order to inform the design of training curricula and to permit deeper longitudinal analysis of career outcomes, the effects of training, and changes in inequality. Senator Ronald Wyden's amendment to the Border Security, Economic Opportunity, and Immigration Modernization Act (S.744, 113th Congress, 2013) was supported by a broad range of unions and associations¹. The Secretary of Labor's congressionally mandated expert advisory group—the Workforce Information Advisory Council²—as well as its predecessor, the Workforce Information Council, produced reports in both 2018 and 2015 that strongly recommended the inclusion of occupations in wage records; the Secretary of Labor responded that was indeed a high priority (3,4). In 2018, the United States Chamber of Commerce convened public and private organizations to report on ways to gather more granular job competency data from employers (5).

¹ <http://www.ifpte.org/downloads/news/manager/307c.pdf>

² <https://www.doleta.gov/wioa/wiac/>

The potential is enormous. If it were possible to combine new computational tools and administrative wage records to generate an automated crosswalk between job titles and occupations, millions of dollars could be saved in labor costs, data processing could be sped up, data could become more consistent, and it might be possible to generate, without a lag, current information about the changing occupational composition of the labor market.

This paper examines the potential to assign occupations to job titles contained in administrative data using automated, machine-learning approaches. Although there has been little research that directly ties firm-level human resource (HR) data on job titles to occupational classifications, there are intellectual foundations for occupational coding that are largely grounded in the survey world. The first foundation is conceptual: to define each occupation. The second is operational: to translate concepts to standardized protocols. The third is statistical: to infer occupations from the information at hand. The fourth pertains to resources: the implementation of such classifications on an extensive scale given the limited resources available. More generally, we contribute to a much larger set of classification problems, which are increasing in salience with the availability of more transaction data. It is important to understand which tools and approaches enable the new, rich, but unstructured data to be used, while minimizing the need for expensive and slow manual classification. For instance, our data also include information on the material, equipment, and supplies that are purchased on sponsored projects as well.

We use a new extraordinarily rich and detailed set of data on transactional HR records of large firms (universities) in a relatively narrowly defined industry (public institutions of higher education) to identify the potential for machine-learning approaches to classify

occupations. This is, to our knowledge, the first large-scale dataset that draws from such HR records across multiple institutions. These data have several advantages. First, the institutions are relatively large and complex, and they use HR systems similar to those of other large and complex organizations in the rest of the economy. Second, the focus on one industry limits the number of possible occupational categories, permitting a targeted analysis. Third, the focus on public universities is attractive because the HR descriptions associated with job titles are available online, and can be used to provide additional information for classification purposes. Finally, the industry is interesting in its own right. Indeed, the production of research often involves the use of intangible assets, particularly labor inputs, and accurate classification of those labor inputs is important for the measurement of scientific productivity.

We build a training dataset from the HR records using human curation and additional rich data sources. First, university staff and trained students manually assign occupations to job titles. That manual curation is then enhanced with additional information from online job descriptions as well as Census Bureau micro-level information on demographic characteristics and earnings. The data are then used to train machine learning models to predict occupations from job titles. Finally, the results are evaluated.

While our results suggest that occupations can be assigned from job titles, they also point to real challenges. In particular, our analysis suggests that there are substantial limits to using machine learning to create discrete occupational categories, even with rich data sources. There are two core problems. The first is that occupational classifications are inherently noisy, so it is difficult to identify ground truth, particularly in a dynamic and changing economy. The second is that job titles have insufficient consistency or detail

across institutions necessary for robust supervised machine learning. We do find that a large number of relatively sparsely populated job titles—a quarter of the titles have only one employee, and over half have fewer than ten employees—could be assigned algorithmically, greatly reducing cost with little impact on accuracy.

II. Background

A major reason for developing occupational classifications is to provide an easy-to-measure pathway from generally understood job activities to skill needs in the economy. The need to capture information on occupations to inform businesses, government agencies, students, and career counsellors about the levels, trends, and changes in skill needs is a continuing theme in national and local workforce policy (6). There are also academic reasons. Occupational classification is deeply rooted in sociology (7), as intrinsic to the measurement of the sources of inequality, social stratification and class mobility. Occupational classification is also essential in economic analyses, describing structural changes caused by technological advancement, automation, globalization, and change in immigration laws (8).

The current approach to occupational classifications is thorough and thoughtful, but quite costly. In addition to cost, the measurement challenges of categorizing worker occupations on surveys are well known: they are notoriously noisy (2). In probably the best known analysis, Mellow and Sider find that only 83.3% of CPS respondents' major (1-digit) occupations match their employer's reports and that share falls to 59.7% for detailed (3-digit) occupations (these rates are considerably lower than those for industry of employment, at 93.1% and 85.4% for major and detailed industry) (1). Bound et al. find similar errors in their overview of measurement errors (9), as do Abraham and

Spletzer (10). Fisher and Houseworth find that there is systematic inflation of occupations for lower-skilled individuals (11).

As noted in the introduction, there has been high-level interest in requiring firms to report occupational data as part of their federal reporting requirements. Both the Workforce Information Advisory Council, an expert group formed to advise the Secretary of Labor, and its predecessor, the Workforce Information Council, recommended adding occupational classifications to unemployment insurance wage records (3,4). In the latter case the group surveyed forty-four states and territories, forty-seven national, state, and regional organizations (representing over 20 million data users in business, education, labor, policy development, economic research, and workforce preparation fields), and rated the need for capturing occupational information as one of the highest of their priorities. The interest in using unemployment insurance wage records for decision-making was also highlighted in the recent report by the Commission on Evidence Based Policy Making (12), and is certainly part of the focus of implementing the legislation that resulted from their recommendations (13).

However, the cost of collecting occupational data manually might well be prohibitive—the state of Texas surveyed businesses and estimated “that the initial cost to employers could range from \$478 million to \$1.2 billion, with annual recurring costs of \$342 million to \$715 million. Costs to the Texas Work Commission were estimated at \$3.1 million in the first year, and a total five-year cost of \$7.9 million to collect this data” (Texas Workforce Commission 2016, p. 17).

Machine learning has become part of the analytical toolkit used by social scientists to automate both classification and prediction tasks (15,16); it “develops algorithms designed to be applied to datasets, with the main areas of focus being prediction (regression), classification, and clustering or grouping tasks” (15). A good overview can be found in the machine learning chapter by Rayid Ghani and Maltz Schierholz in our recent book (17). In the particular context of occupational classifications, there have been attempts to incorporate machine-learning methods by using open-ended survey questions to inform classifications (18, 19). These works found that automated coding was feasible if there is sufficient training data. They emphasized the importance of data preprocessing, algorithmic quality, and thoughtful use of distance metrics in improving occupational prediction. They also suggested that machine learning might also have value by providing responders with candidate occupations as part of a learned cluster, rather than as part of a constructed and hierarchical decision tree. This approach, which is very different from ours, places a higher burden on respondents. In contrast, we use administrative information on job titles, rather than survey responses. We provide more detail in the technical discussion below.

III. Data and Framework

The administrative data we use are derived from the UMETRICS project, which builds on and extends the federal STAR METRICS effort (20). These data are maintained by the Institute for Research on Innovation and Science (IRIS) at the University of Michigan and currently contain record-level information on all wage payments made to individuals

through research grants at 26 participating research universities (20,21). In the interest of homogeneity, for our analysis, we chose large public research universities in the Big 10.³

Although multiple files are provided by the universities, we focus on the employee file, which for each federally funded project, contains all payroll charges for all pay periods (period start date to period end date) with links to both the federal award ID (unique award number) and the internal university ID number (recipient account number). Also available from the payroll records are the employees' internal de-identified employee number, the job title, their Full Time Equivalent (FTE) status, and the proportion of earnings allocated to the award. In addition, the UMETRICS program has incorporated additional fields (notably, the name and date of birth of those supported on federally funded projects) to enable data linkage, and has enhanced the core data with additional information on grants derived from public sources.

We view these data as a valuable laboratory for quantifying the prospects for a machine-learning approach to occupation classification. In some ways these universities are well-suited to a machine-learning approach—they are large, generally similar, and highly structured. Thus, we can identify many different categories of workers and assess our ability to identify similar workers at other institutions. On the other hand, the uniformity of these institutions makes our task somewhat more challenging in that we need to make relatively fine distinctions (e.g. a dataset comprised of longshoremen and financial analysts would have more variability than our data).

³ The universities are Indiana, Wisconsin, Iowa, Michigan, Minnesota, Penn State, Rutgers, and Ohio State University.

In determining occupational classifications, we drew heavily on standard principles. We were particularly interested in building a classification system that described the way people are used in the production of research. Our classification system benefited from extensive consultation with universities, which identified five core characteristics that distinguish personnel employed on research projects: (i) Permanence in their position (ii) Research Role, (iii) Professorial Track, (iv) Scientific Training, and (v) Clinical Association. These core characteristics are similar to ones used in Standard Occupational Classification (SOC) system: classification principle #2 reads “Occupations are classified based on work performed and, in some cases, on the skills, education, and/or training needed to perform the work at a competent level.”

Based on this input, we iteratively developed a hierarchical occupation classification system. In the end, we identified a two-level classification system. The first level is based on a person’s relationship to the university – faculty, undergraduate, graduate student, postdoc, or staff/other. In the second level, we subdivide staff/other based on function. Figure 1 lays out our classification system and Appendix I provides illustrative job titles for the occupations.

As we discuss in detail in the following sections, we manually assigned an occupation from our classification system to job titles from the eight universities. Then we used this manually curated data linking job titles to occupations as a training dataset for a

supervised machine learning approach that algorithmically assigns occupations to job titles.⁴

FIGURE 1 GOES HERE

IV. Creating a Training Dataset from HR Records

The first step was to manually classify occupations based on job titles, which points to the scale of the problem and hence the value of an automated approach. First, the total number of job titles varied from the low hundreds to low thousands across universities—it is likely that similar variation occurs in firms in other sectors of the economy.

The composition of the research personnel by occupation is shown in Table 1.⁵ Also shown in Table 1 is the average number of person-years by occupation for the four largest and four smallest universities (i.e., those universities whose total number of person-year counts is above or below the median). Big universities have, on average, twice as many research personnel paid by research grants, and the share of graduate and undergraduate students is somewhat larger for the big universities.

TABLE 1 GOES HERE

⁴ Our sample consists of individuals appearing in the UMETRICS employee file between 2012 and 2014. Universities that are missing records in any year between 2012 and 2014 were dropped. Universities that had fewer than 100 employees in any occupational class were also dropped because the accuracy of classification algorithms may not be reliably calculated. Eight universities satisfied these sample restrictions.

⁵ The occupations Staff and Others were combined into a single category because the distinction between the two classes is somewhat ambiguous and less important. The unit of observation is a person-year. That is, an individual can be counted up to three times, once per calendar year. Because career transitions can happen within a calendar year (e.g., an individual changing his or her occupation from graduate student to postdoctoral researcher over the summer), only individuals who appeared in the employee table under the same job title both before July 1 and after September 30 were included in our sample.

It is also worth noting that there is substantial variation in the number of people with each job title, as reflected in the average number of people per job title and the fractions of job titles that contain different numbers of people. We divide universities into two groups—the four with the “coarsest” and the four with the most “detailed” job titles. As shown in Table 2, for the universities that use more detailed job titles, as many as 30% of job titles had only one employee. For the universities that use coarse job titles, the proportion occupied by the job titles with more than 100 employees is nontrivial, and job titles with more than 1000 employees were not uncommon. This has important implications for our work—the handling of some job titles has a much greater effect on accuracy of the entire occupation classification than others.

TABLE 2 GOES HERE

Even using the relatively straightforward categorization depicted in Figure 1, we identified three separate measurement challenges that will almost surely be manifested in other firms across the economy. Each results in issues that affect the quality of the training data.

First, when different employees with the same job title perform different tasks, the same job title can map to two distinct occupations. For instance, consider employees with the job title of “program coordinator”. In some cases, these employees may be managing the business operations of a scientific research program at a university center and should thus be assigned the occupation “Research Facilitation Staff”. In other cases, these employees may be involved in educational or student experiences and should thus be assigned the occupation “Instructional Staff”. In this case, different people with the same job title

perform different tasks and should thus be assigned to different occupations. This implies that a full classification must operate at the level of individuals rather than job titles.

Second, some job titles are at the margins of categories. For instance, consider employees with the job title “laboratory supervisor”. In many cases, these employees appeared to perform some tasks that would suggest assigning them the occupation “Research Facilitation Staff” and other tasks that would suggest assigning them the occupation of “Research Staff”. For instance, a laboratory supervisor may serve as an administrator for a university research lab and also conduct research within the lab. Because such employees’ work encompasses the responsibilities of two occupations, it can be argued that they fall at the margin of the occupational categories, which points to the value of a task/skill-based classification versus a categorical classification. This measurement challenge is conceptually distinct from the first insofar as a single individual performs functions that cross categories, rather than two separate people with the same job title performing different functions.⁶

The third measurement challenge is ambiguity: vague titles limited our ability to confidently assign occupations to job titles. “Administrative support”, “coordinator”, and “professional aide” are examples of unclear job titles. Some employees with these titles work in human resources, undergraduate admissions, or a wide range of offices supporting general university functions, while other employees with these titles are

⁶ Although jobs at the margins of categories are not limited to managerial jobs (and our categories are carefully chosen to minimize such uncertainty), managerial jobs often lie at the margins of categories because they require expertise in different kinds of skillsets. One way to address this issue is to create management occupations. For our data, the number of job titles at the margins of categories is relatively small; therefore, we proceeded without creating managerial occupations.

directly involved in supporting or conducting scientific research. To a large extent, this ambiguity reflects a fundamental noisiness in occupational classifications in their own right.⁷ When dealing with ambiguous titles, researchers should be aware that it could influence the learning process of machine-learning algorithms if manually classified occupations were subsequently used for training. For example, the job title “Student help” can belong to either a student who provides help or a staff member who helps students. If we assign this title to a student occupation, we implicitly reinforce the association between the word “student” in the job title and the title belonging to a student occupation, potentially increasing the chance of misclassification for job titles such as “Student learning center coordinator”. Another example of this type is “Fellowship”, which may be intended to mean “fellow”, usually a graduate student, or a staff member who handles administrative work involving fellowship. Addressing title ambiguity is conceptually straightforward, but it requires a great degree of cooperation from data-submitting organizations.

It is worth noting that the same person can have multiple relationships to a university. For instance, a student may hold a staff position or a staff member can become a student to take advantage of a discount on tuition. In this case, the person would be both a staff member and a student. Such multiple relationships pose a challenge, but also present an opportunity for obtaining unique data on career paths. The ideal handling of such cases

⁷ We benefited tremendously from input from member universities that provided extensive input on our classification approach up front, provided a wealth of data, and that have, in many cases, provided extensive feedback on our classification of their employees, especially to address the issues above.

depends on the intended use of the data. If one wants to measure the inputs to a production function, then the preferred approach would likely be to assign the person to the staff title (i.e., to the role that he or she is playing on the sponsored project in question). If the goal is to identify people who have studied at the university, the preferred approach would be to assign the person to the appropriate student occupation. Our data tend to favor the first approach because the primary classification is based on the job title.

Another issue that generates a challenge, but also has the potential to enrich the data greatly, is that people's relationship to a university may change over time. An undergraduate may graduate and enter a graduate program at the same school or take a job as a staff member. A graduate student may take a staff, faculty, or postdoc position upon completion of his or her degree. Obviously, some such pathways are more likely than others. These transitions potentially provide additional leverage on the classification of specific job titles and also provide rich data on career paths.

Incorporating additional external information

We use several different sources of external information, including online job descriptions, publicly available electronic salary databases, university and professional networking websites, and historical administrative earnings and employment data.

Many firms will have HR descriptions that map directly onto job titles. This information could, in principle, provide substantial external information that can be leveraged for occupational classification. In our case, the eight universities had searchable databases

for employment and job postings on university HR websites. These typically provided detailed descriptions of specific job titles to confirm the nature of an employee's work. When these descriptions failed to provide the necessary information to correctly classify a position, electronic salary databases for public universities proved to be particularly helpful sources of information on employee names. Using names and job titles enabled us to examine individual profiles on university and professional networking websites, both of which offered detailed explanations of employees' work. Specific information on actual employees rather than just their titles enabled a more careful classification of similarly related positions in some cases.

Placement and earnings are obtained by linking UMETRICS data to data at the U.S. Census Bureau. Given large differences in age and earnings between various occupations in our data, information on an individual's age and earnings can provide valuable information about that individual's occupation. Employees in the UMETRICS data are linked to Census data using a Protected Identification Key (PIK), Census's internal anonymized individual identifier.

V. Measurement and Standardization

The development of clear standardized protocols for interviewers is critical for consistent measurement across individuals. Similarly, good measurement is critically dependent on developing consistent protocols for preprocessing the data so that measures can be standardized across businesses. This is particularly important since each business will have different shorthand to classify job titles. In this section, we will illustrate the

challenges of standardizing data collected across multiple organizations with different conventions. We will focus on the abbreviated nature of job titles, but we expect similar challenges will arise in processing texts describing job responsibilities, salary grade, retirement benefits, and other information that may be available.

To automate the classification process, we first need to convert job titles to numeric values because most machine-learning algorithms accept only numeric inputs. For short texts like job titles, the most common way of converting texts to numeric features (equivalent of regressors in regression analysis) is to record the presence/absence of keywords. For example, if we have job titles “research analyst” and “research support”, the array of feature names is [“research”, “analyst”, “support”] and the text-to-feature conversion would return the vector [1, 1, 0] for “research analyst” and [1, 0, 1] for “research support”. These vectors will then be used as inputs for machine-learning algorithms that predict occupations.

One problem with this approach is different abbreviations/synonyms in the job titles may represent the same feature. For example, it is clear to humans that “assistant” and “asstnt” both represent “assistant”, but machines treat them as different features. To avoid creating separate features for different abbreviations of the same word, job titles need to be normalized before being converted to numeric vectors.

Because creating a normalization mapping is labor intensive, one may be tempted to use edit distance to determine whether a string of letters is an abbreviation of a word.

However, generic edit distance fails to address challenges that are specific to abbreviations: for instance, both “busin” and “buses” are formed by deleting three letters from “business” and therefore have the same edit distance; however, the former is more likely to be an abbreviation for “business”. Developing a set of rules for determining the validity of abbreviation is not a trivial task. Though the disabbreviation algorithm we developed is imperfect, we employed the algorithm for the subsequent analyses to reduce noise in the data (see appendix I for details).

VI. Machine Learning

We first explored a wide range of classification algorithms, including linear regression. We then selected a few algorithms that seemed to work well for our project and conducted a preliminary analysis comparing their performance. The algorithms that made our “short list” are Multinomial Naïve Bayes, Bernoulli Naïve Bayes, Random Forests, and Extra Trees (Extremely Randomized Trees). We will briefly describe each algorithm below, but interested readers may refer to, for example, James et al. (22) for more details.

The Naïve Bayes classifiers compute the conditional probability of an observation falling in a certain class (equivalent to discrete “y” variable in a regression) given features (equivalent to covariates “x” in a regression) using Bayes rule. The Multinomial Naïve Bayes classifier assumes that the conditional probability that each feature appears given a class follows a multinomial distribution. The Bernoulli Naïve Bayes classifier assumes that the conditional probability of the presence or absence of features given a class follows a Bernoulli distribution. Because it is unlikely that the same word appears more than once in a job title, the Bernoulli Naïve Bayes classifier is more suitable for our

purpose. The major disadvantage of either of the Naïve Bayes classifiers is that they both rely on the underlying assumption that the features are independent. This means, for example, given that the job title belongs to a graduate student, the presence of the word “research” cannot change the probability of also observing the word “assistant” in the job title. Because the assumption of independent features is most likely violated for our case, we rejected Naïve Bayes classifiers.

Random Forest and Extra Trees are both tree-based algorithms. We begin by describing a simple tree algorithm. Figure 2 shows part of a decision tree that classifies employees into the main five classes from Figure 1 (faculty, postgraduate students, graduate students, undergraduate students, and staff/other) based on their job titles. Each box contains (i) branching rule; (ii) Gini impurity; (iii) number of observations contained in the node; (iv) composition of observations; and (v) majority class.

FIGURE 2 GOES HERE

Branching rules specify the feature name and the cutoff value. For example, at the top node, job titles that contain the word “graduate” less than or equal to 0.5 times follow the left branch, while those that contain the word “graduate” more than 0.5 times follow the right branch. Because feature values are integers, it is equivalent to the following: job titles without the word “graduate” follow the left branch and those with the word “graduate” follow the right branch. If “graduate” is not present, it next tests for “professor” and if “graduate” is present, it tests for “post” (as in postgraduate). Note that the node at the bottom right does not have the branching rule because it is a terminal node.

“Samples” represents the number of observations in each node. “Value” lists the number of faculty, postgraduate students, graduate students, undergraduate student, and other in that order. “Class” is the mode of the class in each node.

Finally, “Gini” reports the Gini impurity. Notice that the Gini impurity decreases as one goes down the tree and reaches 0 when a node consists of one class (bottom right node).

It is calculated as follows:

$$G = \sum_{c=1}^5 p_c(1 - p_c),$$

where p_c is the proportion of class c observations at the node.

Although simple and easy to interpret, the Tree algorithm tends to overfit. That is, the algorithm uses too much information that is idiosyncratic to the training set, and thus the predictive accuracy tends to be lower. The Random Forest classifier is intended to mitigate the issue of overfitting by forming a collection of trees. The trees in the forest are slightly different from one another. The variation is generated by introducing randomness to the algorithm. Specifically, each tree is created from different subsample of training data (this is called bagging). Also, when branching, the algorithm does not necessarily choose the feature that minimizes the Gini impurity. The final output (the predicted class) is the class predicted by the greatest number of trees.

The Extra Tree classifier uses the entire training set to create each tree, but introduces randomness by randomly choosing the cut point when branching rather than choosing the

optimal cut point that minimizes the Gini impurity. The random cut point is useful for a continuous feature such as age. For example, it may be that 22 is the optimal cut point for distinguishing undergraduate students from everyone else. However, the Extra Tree may choose a different cut point, say, 20. The Extra Tree can then use other features such as the absence of the word “graduate” to identify undergraduate students who are over 20. Since our feature, the number of times each keyword appears in a job title, is mostly binary (because it is unlikely that the same word appears more than once in a job title), the random cut point would not create much variation: The branching rule “The word ‘research’ appears more than 0.3 times in the job title” is the same as “The word ‘research’ appears more than 0.6 times in the job title”. For this reason, we concluded that there is little gain from using Extra Tree classifier, and decided to use Random Forest classifier.

Our preferred random forest approach (along with the others) is a supervised machine learning algorithm. Thus, it requires a “gold standard” to train the algorithm. Once trained, the algorithm can generate estimates for other samples. In our case, our “gold standard” comes from occupations that have been manually assigned for each university. Although we recognize the potential for human error, we refer to these as the “true” classes. Because we have data on eight institutions, throughout our analysis we estimate our models eight times—holding out data from one university, one at a time, for testing—and use data from the remaining seven universities for training. Thus, for a given set of tuning parameters (discussed below), we grow eight separate random forests, each using

data from one university for testing the accuracy of the forest and using data from the other seven universities for training the random forest.

In our analysis, we used the predictive accuracy as a measure of performance. Formally, the accuracy is defined as

$$Accuracy = \frac{\#(predicted\ class = true\ class)}{\# total\ observations}.$$

Random forests have three main tuning parameters: 1) the total number of features supplied to the random forest, 2) the number of features to be considered at each node of the tree, and 3) the number of trees grown in the forest (i.e., the number of samples randomly selected to build a decision tree). The tradeoff of including more features overall is between having more features to improve prediction and overfitting because of idiosyncratic relationships that may be present in the data. We filter out noise in the sample by pre-selecting the features to avoid overfitting idiosyncratic relationships that may be present in the sample. The total number of features used in the random forest controls the amount of noise to avoid overfitting. The number of features that the random forest can choose between at each stage controls the variability of the trees: the smaller the set of features to be considered, the more variable the trees become because there is more randomness in the selection of the feature. In the extreme case where only one feature is considered at each split, the selection of feature is totally random (i.e., whichever feature is selected becomes the one used for branching).

All analyses were done in Python and the construction of random forests was done using the Scikit-learn package (23).⁸ The package is heavily used by social scientists because, as the authors note, it is a Python module that integrates “a wide range of state-of-the-art machine learning algorithms for medium-scale supervised and unsupervised problems. This package focuses on bringing machine learning to non-specialists using a general-purpose high-level language. Emphasis is put on ease of use, performance, documentation, and API consistency. It has minimal dependencies and is distributed under the simplified BSD license, encouraging its use in both academic and commercial settings. Source code, binaries, and documentation can be downloaded from <http://scikit-learn.sourceforge.net>” (p2826 (23)).

The RandomForestClassifier module of the Scikit-learn package allows the users to change the parameters mentioned above. To determine the total number of features supplied to the random forest, we fit a decision tree, for each training set, using all 1-grams and 2-grams that appeared in the job titles. Then, the feature importance score was calculated, and the features with the highest importance scores were selected, varying the score cutoff. The total number of features that were fed into the model varied depending on which university was reserved for testing, but was roughly 50, 100, 200, 500, and 7000, where 7000 is the total number of 1-grams and 2-grams appearing in the job titles in the training set and 500 is the number of features that had a strictly positive importance

⁸ The work was performed at the UCLA Federal Statistical Research Data Center (RDC). The RDC compute nodes were IBM HS22V blade servers. Each had 12 CPU cores (24 with hyper-threading enabled) and 288GB of memory. The version of Scikit-learn that is used in the Census Bureau’s research1 machine is 0.18.1; we believe that is the same version that was used for this work.

score. We also varied the number of features considered at each split (default is the square root of the total number of features supplied to the random forest). Finally, we varied the number of trees grown in the forest, in increments of 100, between 100 and 1,000.

In determining the optimal parameter setting, we considered both unweighted and weighted accuracy. The unweighted accuracy was computed treating each job title as one observation – no matter how many employees have that job title, the title receives a weight of 1. The weighted accuracy was computed treating each individual as one observation; equivalently, job titles were assigned a weight equal to the number of employees that have that job title. The most important tuning parameter for determining classification accuracy was the total number of features provided to the random forest (which is implicitly determined by the importance score cutoff). The fraction of features to be considered at each node and the number of trees grown had a minimal effect on the accuracy. Based on the overall weighted and unweighted accuracy, the optimal parameter setting limits the number of features supplied to the random forest to about 200 and uses the default setting of the square root of the total number of features to be considered at each node.

FIGURE 3 GOES HERE

Figure 3 shows the accuracy (proportion of correct prediction) for each level of predicted probability (the probability share of the predicted occupation indicated by the posterior

distribution returned by the algorithm). The overall accuracy varied from 60% to nearly 100% regardless of whether the data are weighted by number of job titles or individuals⁹. Although random forests can potentially increase the efficiency of occupational classification, an average accuracy of about 80% may not be high enough to justify a total replacement of manual classification by automated machine-learning algorithms. These results reinforce our belief that the predicted probability and the number of individuals that hold a job title should be jointly used to identify job titles for manual review.

We see two (potentially complementary) roles for machine learning in occupation coding and other similar bucketing tasks. One approach is to use an algorithmic approach to classify uncommon job titles. Such cases are (by construction) plentiful and have a relatively small effect on the overall accuracy of the classification. The second role is to accept only predictions with concentrated probability mass at one class. In other words, to adopt the prediction only when the random forest classifier is “confident”. Obviously, these two approaches could be combined—defining “isoquants” over the size and accuracy to trigger manual review. In this approach only relatively large, uncertain job titles would be reviewed manually.

Robustness

We explored a wide range of modifications of our basic approach to try to obtain performance improvements. Here, we outline the analyses we performed and their main results. Appendix II provides details on both the analyses and their results.

⁹ Unweighted accuracy is the proportion of job titles whose predicted class matched the true class. For weighted accuracy, the number of employees for the job title is used as a weight.

For the eight universities used in the above analyses, the number of employees ranged roughly from 5,000 to 20,000. When we train the random forest classifier on seven universities, it is possible that the shape of a tree is heavily influenced by a few universities in the training set with a large number of employees. To investigate this possibility, the training set was modified so that universities in the training set have roughly equal numbers of employees. The modifications were made in two ways: inflating and deflating. As Section A in Appendix II shows, there was no significant change in the accuracy with these modifications.

The number of employees per job title ranged from 1 to nearly 10,000 for the eight universities, with the average being 24.4 employees per title. Concerned that the titles in the training set are “too noisy”, we investigated the effect of dropping thin titles (varying the threshold at which a title is flagged as “thin” from 5 to 50 employees) from the training set. We recorded the average predictive accuracy for titles with different numbers of employees. Again, there was no significant change in the accuracy with these modifications. These results are discussed in Section B of Appendix II.

We observed that some titles that could be easily classified manually like “Graduate Assistant” are not always correctly classified by our random forest. This appears to be caused by the existence of “extraneous” information in some job titles. To address this issue, we applied partially unsupervised learning. In particular, titles that (after applying the job cleaning algorithm outlined in the appendix) contain the words “faculty”,

“professor”, “postgraduate”, “graduate” or “undergraduate” were classified first and then the random forest classifier was applied to the remaining titles (both the training set and the test set consist of titles that do not contain any of the words listed above). The effect of this partially unsupervised learning on the predictive accuracy is small, with our classification for some universities improving and others degrading. See Section C of Appendix II for details.

Census Bureau links permitted us to examine whether having information on individuals’ age and earnings increased the quality of prediction. These variables would appear to be valuable predictors, especially in this context because of the large differences in ages and earning across occupations. As shown in Table A4, there is some gain, but it is not extraordinarily high across the board. The largest gains, by far, are for undergraduates when occupations are weighted by the number of people in them.¹⁰

As indicated in the previous sections, people can hold multiple titles at a point in time (or in close succession) and can transition between titles. As some transitions are more common than others (i.e., transitions from undergraduate to graduate and/or from graduate to postgraduate and/or from postgraduate to faculty are more common than the reverse transitions), it is possible to use transitions between titles and concurrent titles (more precisely, occupational classes that are associated with these titles) as predictors in the random forest classifier to improve predictive accuracy. Transitional and concurrent

¹⁰ Because there are considerably more staff members than undergraduates overall, there is a tendency for the random forest to misclassify undergraduates as staff.

titles can also be used to identify unlikely transitions in the “ground-truth” data, providing an opportunity for a revision. Beyond improving the accuracy of the data, exploring concurrent positions and transitions can add to the richness of our data by providing information on career paths. Section D of Appendix II provides details of this analysis.

Using concurrent job titles and the transitions between job titles involves some form of iterative procedure. Appendix II details a number of issues related to using transitions and concurrent titles. As a first step toward incorporating transitional and concurrent classes into the random forest classifiers, we included the manually classified transitional and concurrent classes in our training data in the model rather than predicted occupations. The resulting predictive accuracy is expected to provide an upper bound for the accuracy obtained from the iterated procedure described in Appendix II. Overall, the use of concurrent titles and transitions across titles has little effect on overall accuracy. In our analysis, no university exhibited a clear pattern on the effect of including transitional/concurrent class as predictors.

Limitation of Machine-Learning Algorithms

Laying aside the issue of developing a classification system, we have discussed four challenges to manual classification. Beyond these issues associated with manually classifying occupations, comparing the predictions made by the random forest and the true class pointed to two possible causes of misclassification. One is unavoidable misclassification, which results from variation in the training data. The other is avoidable

misclassification, which results from the inherent limitations of the random forest classifier.

The first type of misclassification is unavoidable because it arises from the limits of manual classification already discussed, such as job titles that have multiple classifications over universities. This type of inaccuracy cannot be overcome by any classifier: resolution of misclassification requires familiarity with job titling convention at each university. It should also be noted that modifiers can change the classification of a job title within a university. For example, “director” and “associate director” may not belong to the same category within a university.

The second type of misclassification is avoidable. Avoidable misclassifications are due to the limitations of the random forest classifier. Below are examples of misclassified job titles along with the prediction made by the random forest, followed by the true class in parentheses.

- Undergraduate fellow → graduate (undergraduate)
- Temporary visiting faculty → staff/other (faculty)
- Teaching assistant → staff/other (graduate)
- Summer term ra (w/o tuit ben) → staff/other (graduate)
- GR AST ½ → staff/other (graduate)

The first two examples illustrate the tendency of the random forest classifier to rely too much on certain words. The word “fellow” is strongly associated with graduate student. Thus, if “fellow” is selected as a branching rule before “undergraduate”, the job title

“undergraduate fellow” will be buried in a node that is predominantly graduate students. Similarly, the word “temporary” is often associated with a staff member and almost never used for faculty. The partially supervised machine learning algorithm described in the previous section is intended to address these issues.

The third example illustrates failure to utilize very informative words or phrases. The presence of the phrase “teaching assistant” in a job title is a good indicator of the employee being a graduate student. However, the absence of the phrase “teaching assistant” in the job title is not a good indicator of the employee not being a graduate student (i.e., there are many graduate students who are not teaching assistants). Thus, when the phrase “teaching assistant” is used for branching, the resulting decrease in the impurity of the succeeding node is negligible. Since the random forest classifier selects the feature that minimizes the weighted average of impurities at succeeding nodes, the phrase “teaching assistant” is unlikely to be selected.

The last two examples illustrate inability of the random forest classifier to use outside knowledge. A human classifier can infer “w/o tuit ben” means “without tuition benefit” and conclude that the job title is associated with a student. Similarly, “1/2” suggests that the person has a half-time appointment, and therefore is likely to be a student. Thus, one may infer that “gr ast” means “graduate assistant”. As seen in the previous example, these phrases are extremely informative; however, because of their rare occurrence and applicability to only a small fraction of employees, these pieces of information tend to be overlooked by the random forest classifier.

In theory, the misclassifications described above might be reduced by providing more training data, adjusting parameters, appealing to other machine-learning algorithms, or reverting to manual classification.

VII. Conclusions

This paper used a rich dataset—to our knowledge, the first dataset with detailed job titles drawn from HR systems from multiple organizations, combined with job descriptions and information about the characteristics of workers—to examine the potential to use machine-learning techniques for occupational classification. We followed the same conceptual framework as that applied by survey methodologists: to define each occupation, to translate concepts to standardized protocols, and to build an approach that would infer occupations from the information at hand. Even though the data were drawn from very similar organizations, with very similar production functions, we found that machine-learning approaches were not substantially better than manual classifications.

However, we do see the analysis as showing real promise for identifying occupations from job titles combined with a machine-learning approach. The most promising use of the machine learning is that it is an inexpensive way of assigning occupations for job titles that have relatively few people in them and/or for which the algorithm imputes a high degree of accuracy. Because many job titles have only a few people in them, this approach could yield substantial cost savings (almost 80% of job titles have 10 or fewer people). At the same time, an entirely algorithmic approach would be unwarranted in our case.

We also believe that a deeper text analysis of the job descriptions associated with job titles might prove to be a promising approach. Job descriptions typically include information about necessary experience, skills, and education, which is not only of interest in its own right but could be very useful for classification purposes.

We note that the focus on universities as a subject of analysis has weaknesses and strengths. Major research universities are very large and complicated institutions. There may be other industries in which it might be easier to apply machine learning to job titles. At the same time, the institutions in our sample all come from one narrow sector of the economy; they are relatively homogeneous and the data are based on a very specific set of activities (research). We speculate that any classification system for the broader economy would have to be specific to an individual sector or set of sectors.

VIII. Acknowledgements

Disclaimer: Any opinions and conclusions expressed herein are those of the authors and do not necessarily represent the views of the U.S. Census Bureau. All results have been reviewed to ensure that no confidential information is disclosed. This research was supported by the National Center for Science and Engineering Statistics. NSF SciSIP Awards 1064220 and 1262447; NSF Education and Human Resources DGE Awards 1348691, 1547507, 1348701, 1535399, 1535370; NSF NCSES award 1423706; and the Ewing Marion Kaufman and Alfred P. Sloan Foundations. Lane was supported through an Intergovernment Personnel Act assignment to the US Census Bureau. Weinberg is grateful for support from R24 AG048059, R24 HD058484, UL1 TR000090, the National Institute on Aging and the Office of Behavioral and Social Science Research via P01AG039347 (on which Weinberg and his work were supported directly by the National Bureau of Economic Research and indirectly by Ohio State). The research agenda draws on work with many coauthors, but especially Jason Owen Smith. We also acknowledge the tremendous contributions made to occupation coding by Gus Bradley, Larkin Cleland, Luke Fesko, Michael Frank, Dinushi Kulasekere, Tristan Myers, Rachel Notestine, Eric Spurlino, Samantha Stelnicki, Tiger Sun, Varuni Sureddy, Jacob Zeiher, and especially Cameron Conrad.

IX. References

1. Mellow W, Sider H. Accuracy of response in labor market surveys: Evidence and implications. *J Labor Econ.* JSTOR; 1983;331–44.
2. Mathiowetz NA. Errors in reports of occupation. *Public Opin Q.* JSTOR; 1992;352–5.
3. Workforce Information Council Administrative Wage Record Enhancement Study Group. Enhancing Unemployment Insurance Wage Records Potential Benefits, Barriers, and Opportunities [Internet]. Washington DC.; 2015. Available from: <https://www.bls.gov/advisory/bloc/enhancing-unemployment-insurance-wage-records.pdf>
4. Workforce Information Advisory Council. RECOMMENDATIONS TO IMPROVE THE NATION’S WORKFORCE AND LABOR MARKET INFORMATION SYSTEM [Internet]. Washington DC; 2018. Available from: https://www.doleta.gov/wioa/wiac/docs/WIAC_Recommendations_Report_2018-01-25_Final_and_Signed.pdf
5. United States Chamber of Commerce. Developing an Open, Public-Private Data Infrastructure for the Talent Marketplace: Phase 1 Report [Internet]. Washington DC; 2018. Available from: https://www.uschamberfoundation.org/sites/default/files/T3Phase1_Report_FINAL_0.pdf
6. Acemoglu D, Akcigit U, Bloom N, Kerr WR. Innovation, reallocation and growth. National Bureau of Economic Research; 2013.
7. Weeden KA, Grusky DB. The Case for a New Class Map 1. *Am J Sociol.* The

- University of Chicago Press; 2005;111(1):141–212.
8. Alonso-Villar O, Del Rio C, Gradín C. The extent of occupational segregation in the United States: Differences by race, ethnicity, and gender. *Ind relations a J Econ Soc*. Wiley Online Library; 2012;51(2):179–212.
 9. Bound J, Brown C, Mathiowetz N. Chapter 59 – Measurement Error in Survey Data. In: *Handbook of Econometrics* [Internet]. 2001. p. 3705–843. Available from: <http://www.sciencedirect.com/science/article/pii/S1573441201050127>
 10. Abraham KG, Spletzer JR. New evidence on the returns to job skills. *Am Econ Rev*. JSTOR; 2009;99(2):52–7.
 11. Fisher JD, Houseworth C. Occupation Inflation in the Current Population Survey. US Census Bur Cent Econ Stud Pap No CES-WP-12-26. 2012;
 12. Commission on Evidence based Policy. The Promise of Evidence-Based Policymaking [Internet]. Washington DC; 2017. Available from: www.cep.gov
 13. Hart N, Shaw T. Congress Provides New Foundation for Evidence-Based Policymaking [Internet]. Washington DC: Bipartisan Policy Center; 2018. Available from: <https://bipartisanpolicy.org/blog/congress-provides-new-foundation-for-evidence-based-policymaking/>
 14. Texas Workforce Commission. Report to the Sunset Advisory Commission Study on the Collection of Occupational Data. Austin, Texas; 2016.
 15. Athey S. The impact of machine learning on economics. In: *The Economics of Artificial Intelligence: An Agenda*. University of Chicago Press; 2018.
 16. Kleinberg J, Lakkaraju H, Leskovec J, Ludwig J, Mullainathan S. Human decisions and machine predictions. *Q J Econ*. Oxford University Press;

- 2017;133(1):237–93.
17. Ghani R, Schierholz M. Machine Learning. In: Foster I, Ghani R, Jarmin R, Kreuter F, Lane J, editors. *Big Data and Social Science: A Practical Guide to Methods and Tools*. Taylor & Francis; 2016.
 18. Bethmann A, Schierholz M, Wenzig K, Zielonka M. *Automatic Coding of Occupations*. 2014;
 19. Tijdens KG. Reviewing the measurement and comparison of occupations across Europe. *AIAS Work Pap. Citeseer*; 2014;(149).
 20. Lane J, Owen-Smith J, Rosen R, Weinberg B. New linked data on science investments, the scientific workforce and the economic and scientific results of science. *Research Policy*. 2014.
 21. Weinberg BA, Owen-Smith J, Rosen RF, Schwarz L, Allen BM, Weiss RE, et al. Science Funding and Short-Term Economic Activity. *Science* (80-) [Internet]. 2014 Apr 4;344(6179):41–3. Available from:
<http://www.sciencemag.org/content/344/6179/41.short> and available in full at
<http://www.cssip.org/work>
 22. James G, Witten D, Hastie T, Tibshirani R. *An introduction to statistical learning* (Vol. 103). New York: Springer. MATH; 2013.
 23. Pedregosa F, Varoquaux G, Gramfort A, Michel V, Thirion B, Grisel O, et al. Scikit-learn: Machine learning in Python. *J Mach Learn Res*. 2011;12(Oct):2825–30.

X. Tables

Table 1. Number of employees paid by research grant by occupation.

Occupation	All universities Total	Big universities Average	Small universities Average
Faculty	16,000	2,600	1,500
Graduate	17,000	3,100	1,200
Staff / Other	29,000	4,700	2,600
Postdoc	6,900	1,100	650
Undergrad	9,700	2,000	450
Total	79,000	13,000	6,400

Note. The table shows the number of employees paid by research grants at all universities in our data and those with more than and fewer than the median number of person-year pairs. Numbers are rounded for disclosure protection reasons.

Table 2. Variation in the volume and size of job titles across universities.

	All universities	Universities with coarse job titles	Universities with detailed job titles
Total number of job titles (across universities)	3,200	1,100	2,200
Total number of employees (across universities)	79,000	48,000	31,000
Average # employees per title (at each university)	24.4	44.4	14.5
1 employee	25%	16%	30%
2-10 employees	54%	52%	55%
11-100 employees	17%	26%	13%
>100 employees	4%	7%	2%

Note. The table shows the distribution of the number of employees per title at all universities in our data, the four universities with the smallest, and the four universities with the largest numbers of employees per job title.

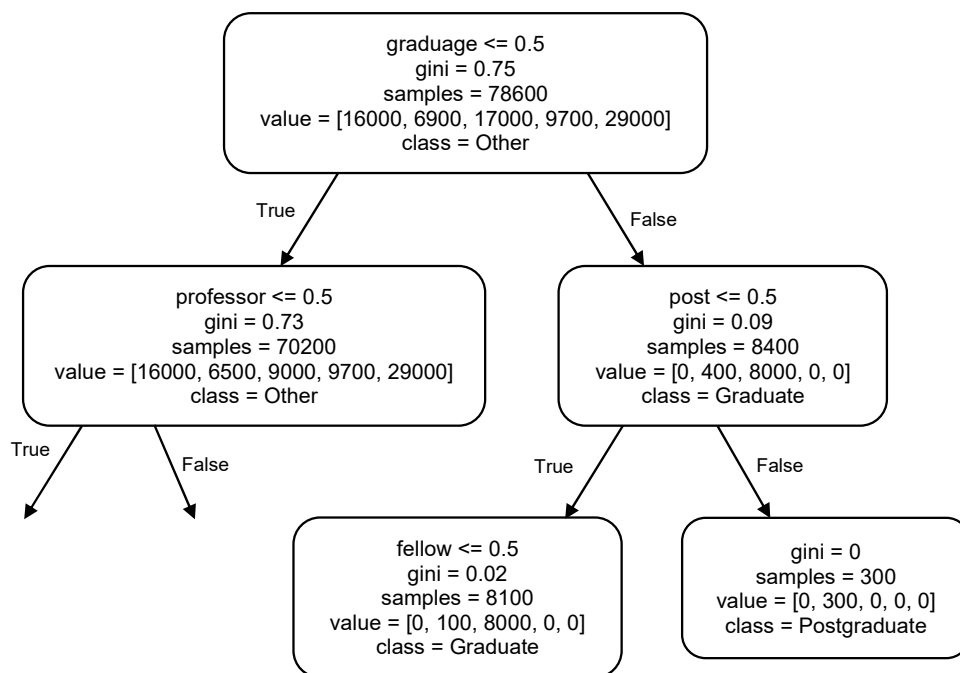
XI. Figure Captions

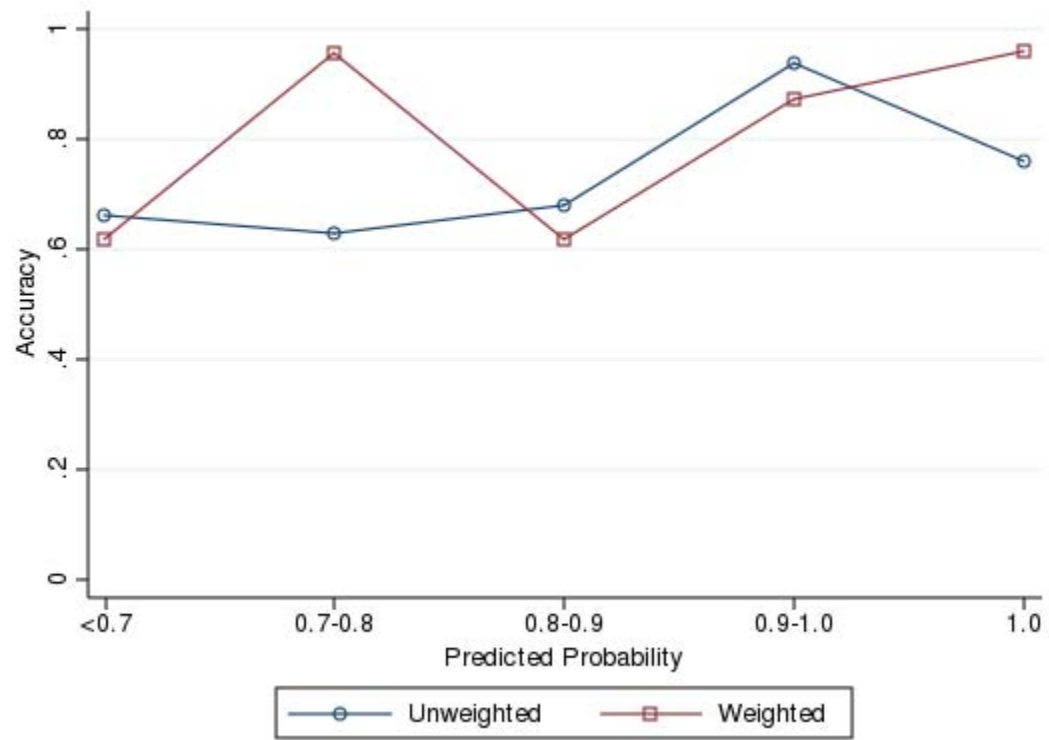
Figure 1. Classification System.

Figure 2. Example of a decision tree. Each node has a keyword indicated at the top of the box. All observations that have the keyword in their job titles follow the right branch, while observations without the keyword follow the left branch. When an observation reaches a terminal node like the one at the bottom right, the class of the node becomes the predicted class for the observation.

Figure 3: Classification Accuracy relative to Predicted Probability. The figure shows the probability that an occupation was correctly coded as a function of the probability that the algorithm predicts it was correctly coded. The unweighted series treats job titles as the unit of observation. The weighted series treats individuals as the unit of observation.

XII. Figures





Appendix I

1. Detailed Description of Occupations

This section lays out the occupation categories that we use, their conceptual definition, and some illustrative job titles. The aggregate occupations are listed first. Staff are subdivided into additional categories, which are laid out below.

1. Faculty

All advanced academic employees who are directly involved in scientific research and/or scientific instruction. These included Deans, Provosts, Tenure/tenure track, Clinical, Research, Visiting Professors, Academic specialists, Center directors.

2. Post Graduate Research

All individuals holding terminal degrees (PhD, MD) who are in temporary training status. These included Postdoctoral, Medical residents/interns/fellows, Clinical fellowships, Research Associates (depends on the university).

3. Graduate Student

Students earning advanced degrees: Graduate students (part time, full time), Medical/dental/nursing/students, Research Assistants.

4. Undergraduate

Students earning baccalaureate/other degrees including full time, part time, summer research assistants, work study; includes high school students who would likely be acting

in a similar capacity. These included Undergraduate students, High school students, Interns/student workers, Nursing students in BA programs.

5. Staff / Other (Not Elsewhere Classified)

Positions that support general university functions such as undergraduate education and student activities. Employees whose titles cannot be attributed to the scientific research enterprise. These included at the aggregate level: **Staff** Instructional, Research, Research Facilitation, Technician, Clinical, Other Staff. The disaggregated staff categories include the following:

5.1 Clinical Staff: All non-faculty health care professionals, Nurses (non-faculty),

Dieticians (non-faculty), Nutritionists, Social workers, Physical therapists, Clinical psychologists, Dental hygienists, Genetics counselors.

5.2 Instructional Academic Specialists: Lecturers, Instructors, Adjunct Professors.

5.3 Research Facilitation: Non-faculty, high level administrators – asst. dean/asst.

provost, associate or assistant center director, Operations managers/managing directors, Administrative/clerical staff – any kind, Finance staff, Regulatory staff, Clinical or clinical research support staff, Laboratory aide, Data collection/interviewer, Media jobs: Graphics/writer/editor/communications, Grants management & administration, Individuals who serve as managers/coordinators/facilitators for laboratory studies/clinical trials/large facilities/research programs (they direct and influence scientific research activity from the level of the laboratory up to the level of the university/research center), Research

dean/provost/administrator, Facility director/administrator, Clinical research administrator, Study coordinators, IACUC coordinators, Clinical trials/research coordinator, Project/Program manager/coordinator, Lab coordinator (not lab manager), Facility/repository manager/coordinator.

5.4 Research Staff: Work likely focuses on scientific aspects of research. All advanced degree qualified, non-faculty scientists and engineers; Research specialist/engineer: Work likely focuses on advanced research analysis; Research professional/specialist; Statistician, bioinformaticist; Research associate (depends on the university); Skilled and specialized employees who have been specifically trained in some area of science & technology; Science Technicians: All technical staff including animal technicians, machinists, mechanics (the category usually includes some reference to a research facility along with the title ‘technician’); Lab manager ; Medical or clinical technician; Research data technician; Regulatory officer (environmental, chemical safety, industrial hygienist); Technical engineer.

5.5 Technician: Administrative and technical employees who are not specifically employed for scientific research purposes but perform job tasks that support the research enterprise; Information technology managers & staff; Software engineer; Data entry/data analyst; Network and systems support.

5.6 Staff Other All other research staff that do not clearly fall into another category.

2. Normalization

We developed a rule-based job title cleaning algorithm. In particular, we created a mapping from abbreviation to normalized word. For example, “grad” is mapped to

“graduate” and “mngr” is mapped to “manager”. The list of abbreviations and possible normalized words were obtained from job titles from eight universities in the UMETRICS dataset, and mappings were created manually.

Abbreviations with multiple possible normalizations were noted (e.g., “res” can be an abbreviation for “research” or “respiratory”; “ast” can be an abbreviation for “assistant” or “astronomy”). Then context-specific normalization (i.e., normalization of phrases) was attempted. For example, both “res” and “ast” are ambiguous abbreviations; however, when they are combined, one can infer “res ast” is an abbreviation for “research assistant”. Normalizing rules for phrases were manually generalized using regular expressions.

When an abbreviation could represent either a person or a field (or an object) that are closely related, we chose the field in general. For example, “scien”, “enginee”, and “crimnl” were normalized to science, engineering, and criminology, instead of scientist, engineer, and criminologist, respectively. The reason is that it seems more harmful to label non-engineers in engineering departments an “engineer” than to label an engineer “engineering”. When an abbreviation is strongly associated with an occupation, however, we normalized it to represent a person. For example, “lect” and “consul” were normalized to lecturer and consultant instead of lecture and consulting, respectively. These are somewhat ad-hoc rules, but these abbreviations are few, and we expect they have a negligible effect on the performance of machine learning algorithms.

When creating the normalization mapping, we preserved common acronyms such as “CSE” for Computer Science and Engineering and “MRI” for magnetic resonance imaging. We expect normalizing these terms has a minimal impact on the predictive accuracy because they identify the fields employees work in but contain little information on tasks they perform.

At the same time the mapping was created, omissions of spaces were noted and a decomposition mapping was created. For example, we encountered job titles such as “rsrchanalyst”, which was added to the decomposition mapping along with the correction “rsrch analyst”. Common stems in compounds, such as bio in biochemistry and neuro in neurosurgery, were not decomposed and compounds were treated like words.

Finally, on the normalization list, we had some abbreviations that are only two letters long. For example, we left “IT” as it is, assuming that it represents Information Technology. However, these could be an abbreviation of some other words or phrases. In our data, we did not find any instances where there was a more suitable normalization, but researchers should be aware that too much guessing when standardizing could introduce more noise than it eliminates.

Aside from working out the details, the major problem with the above described normalization algorithm is that the mapping is not comprehensive. For example, “research” may be mapped from “resear”, “rsrch”, and “resch”, but if there is no mapping from “resech” to “research”, “resech” will remain abbreviated. By comparing manually

normalized job titles and normalization returned by the algorithm, we identified normalizations that were not captured by the normalization mapping, and iteratively revised our normalization mapping. We also wrote regular expressions to normalize words that frequently appear in our data such as “research”, “postdoctoral”, and “administrator”.

3. Coding decisions

There are also some methodological issues of interest. First, we designed our classification to increase certainty: grouping workers whose jobs were so similar that it would be hard to separate them based on job titles (and for whom the value of distinguishing occupations has the least value). Second, we employed a two-level system, where the first-level occupation can frequently be assigned with a high degree of certainty, and much of the uncertainty appears at the second level. Third, we assigned up to two occupations to each job title to allow researchers to probe the sensitivity of results. Fourth, we rated job titles based on the degree of certainty that they were correctly classified on a scale of 1-5. Our coding system was:

(5) The job title serves as an immediate identifier into this classification category or, through research, it is almost certain that it belongs in this category: e.g. Post Doctoral Researcher; Computer Technician.

(4) The job title probably belongs in the category indicated, as supplemented by research on university website.

(3) The job title belongs in the category (either aggregate or disaggregate) with moderate certainty (either very indicative job title or research result, but not both).

(2) The job title is vague and/or ambiguous, but there is some indication that the position belongs in this category.

(1) The job title may belong in this category, but there is little certainty, and the classification cannot be verified through research.

After manual classification, universities were given the opportunity to review and comment on the classification, with their attention drawn to the largest and most ambiguous titles.

Appendix II

A. Different Numbers of Employees

For the eight universities used in the above analyses, the number of employees ranged roughly from 5,000 to 20,000. When we train the random forest classifier on seven universities, it is possible that the shape of a tree is heavily influenced by a few universities in the training set with a large number of employees. To investigate this possibility, the training set was modified so that universities in the training set have roughly equal numbers of employees. The modifications were made in two ways: inflating and deflating.

(1) Inflating

Let

$N_{u,t}$ = number of employees at university u for job title t , and

N_u = number of employees at university u .

Then the modified number of employees is

$$\tilde{N}_{u,t} = N_{u,t} \times \frac{\max\{N_v\}}{N_u},$$

rounded to the nearest integer. For example, if university X has a total of 16,000 employees and if the largest university in the training set has a total of 20,000 employees, the number of employees for each title at university X is multiplied by 1.25 and rounded to the nearest integer. If a title has 3 employees, the inflated number of employees is $1.25 \times 3 = 3.75$, so it will be rounded to 4.

(2) Deflating

Instead of scaling up the number of employees to the level of the largest university in the training set, deflating scales down the number of employees to the level of the smallest university:

$$\tilde{N}_{u,t} = N_{u,t} \times \frac{\min\{N_v\}}{N_u}.$$

For example, if university X has a total of 20,000 employees and if the smallest university in the training set has a total of 5,000 employees, the number of employees for each title at university X is multiplied by 0.25 and rounded to the nearest integer. If a title has 10 employees, the deflated number of employees is $10 \times 0.25 = 2.5$, so it will be rounded to 3. If a title has 1 employee, the deflated number of employees is $1 \times 0.25 = 0.25$, so it will be rounded to 0. In other words, the title will be dropped from the training set.

Results

As evident in Table A1, inflating and deflating the number of employees in the training set has no meaningful effect on the unweighted accuracy. There is a little improvement in the weighted accuracy for big universities when the number of employees in the training set is deflated. One possible explanation is that deflating reduces the noise in the training data because uncommon job titles are dropped from the training set due to rounding if the deflated number of employees is less than 0.5.

Table A1. Accuracy when total weight is balanced across universities

Size of university	Weight	Benchmark	Inflating	Deflating
All universities	unweighted	0.83	0.83	0.82
Big universities	unweighted	0.87	0.86	0.86
Small universities	unweighted	0.80	0.80	0.79
All universities	weighted	0.84	0.82	0.85
Big universities	weighted	0.83	0.82	0.86
Small universities	weighted	0.84	0.82	0.82

B. Discarding Thin Titles

The number of employees per job title ranged from 1 to nearly 10,000 for the eight universities, with the average being 24.4 employees per title. Concerned that the sparsely populated titles in the training set are particularly “noisy”, we investigated the effect of dropping thin titles from the training set. The question we tried to answer is “Do thin

titles negatively affect the learning and consequently degrade the performance of predicting for heavily populated titles?”

For each university, we used the remaining seven universities for training and discarded the titles with fewer than a certain number of employees in them from the training set. We used the threshold of 5, 25, and 50 employees per title. Then we recorded the average predictive accuracy for titles grouped by the number of employees per title: 1–4 employees, 5–24 employees, 25–49 employees, 50–99 employees, 100–499 employees, 500–999 employees, and 1000+ employees. The idea is that dropping titles that have fewer than a certain number of employees from the training set may have different effects on the prediction accuracy for thin titles and for heavily populated titles.

Results

The resulting prediction accuracies are shown in Table A2. Cutoff = 0 corresponds to the benchmark, where all job titles are included in the training set. As the cutoff value increases, more and more job titles are excluded from the training data. There is a small decrease in accuracy caused by discarding uncommon titles for job titles that are relatively thin. In contrast, the accuracy improves as the cutoff increases for job titles with 1000 or more employees. This is expected because highly populated job titles tend to have simple, straightforward descriptions and therefore do not benefit from infrequently used features brought to the training set by uncommon job titles. Indeed, excluding uncommon job titles from the training set makes the training set less noisy,

allowing the random forest classifier to construct better decision trees with a high predictive accuracy.

Table A2. Accuracy when thin titles are discarded from the training set

Size of Title	Weight	Cutoff = 0	Cutoff = 5	Cutoff = 25	Cutoff = 50
<5	unweighted	0.81	0.81	0.82	0.80
5-24	unweighted	0.84	0.84	0.84	0.82
25-49	unweighted	0.91	0.91	0.91	0.89
50-99	unweighted	0.88	0.87	0.86	0.84
100-499	unweighted	0.82	0.82	0.85	0.82
500-999	unweighted	0.88	0.88	0.88	0.88
>=1000	unweighted	0.75	0.83	0.92	0.92
<5	weighted	0.82	0.82	0.83	0.81
5-24	weighted	0.83	0.83	0.83	0.81
25-49	weighted	0.92	0.91	0.92	0.89
50-99	weighted	0.88	0.87	0.86	0.84
100-499	weighted	0.80	0.79	0.82	0.80
500-999	weighted	0.90	0.90	0.90	0.90
>=1000	weighted	0.79	0.85	0.93	0.93

C. Partially Unsupervised Learning

After observing that titles like “Graduate Assistant” are not always correctly classified, we applied partially unsupervised learning. The incorrect classification appears to happen because of “extraneous” information in some job titles. In particular, titles that contain the word (after applying the job cleaning algorithm) “faculty”, “professor”, “postgraduate”, “postdoctoral”, “graduate” or “undergraduate” were classified first and then the random forest was applied to the remaining titles (both the training set and the test set consist of titles that do not contain any of the words listed above).

The resulting accuracies are shown in Table A3. There is no real difference in the unweighted accuracy between the supervised learning benchmark and partially unsupervised learning. Contrary to our expectation, the weighted accuracy deteriorated for the universities with granular job titles, while it improved for the universities with coarse job titles. This suggests a possibility of overfitting; in the absence of very important features such as “faculty” and “undergraduate”, less important features appear to be more important than they actually are. One possible solution is to recalibrate the parameters to filter out marginally informative features.

Table A3. Accuracy using Partially Supervised Learning

University	weight	benchmark	partially unsupervised
all universities	unweighted	0.83	0.83
universities with coarse job titles	unweighted	0.90	0.91
universities with granular job titles	unweighted	0.79	0.79

all universities	weighted	0.83	0.84
universities with coarse job titles	weighted	0.82	0.86
universities with granular job titles	weighted	0.84	0.81

D. Using Age and Wage Data

Census Bureau links permitted us to examine whether or not having information on individuals' age and earnings increased the quality of prediction. These variables would appear to be valuable predictors, especially in this context because of the large differences in ages and earning across occupations. As shown in Table A4, there is some gain, but it is not extraordinarily high across the board. The largest gains, by far, are for undergraduates when occupations are weighted by the number of people in them. The benchmark analysis shows the predictive accuracy for all individuals whose true occupation falls in the occupation indicated in the row heading. The column headed "age and wage" shows the predictive accuracy for individuals for whom we have age and wage information (i.e., subset of the benchmark population). For this subset of population, the "age" column shows the accuracy when age is used along with job title for prediction; the "wage" column shows the accuracy when wage is used along with job title for prediction; the "age and wage" column shows the accuracy when both age and wage are used along with job titles for prediction.

One interesting observation is that the predictive accuracy increases drastically for graduate students with age and wage information. This may be because common jobs titles like teaching assistant are associated with more standardized hiring procedures that increase the chance of students' information being stored in a more organized way on the

university system. This effect, however, disappears when the job titles are not weighted by the number of employees associated with that job title (the bottom half of the table).

This could be due to idiosyncratic job titles and associated non-standardized hiring processes.

The table also shows that correctly classifying undergraduate students is particularly difficult even with the information on age and wage. This is probably due to heterogeneity within the undergraduate researcher body: there are traditional students straight out of high school as well as adult students, who are often employed.

Table A4. Accuracy using age and wage data

Fraction of individuals whose predicted class matches the true class by true class					
Actual Occupation	Benchmark	Sample with Age & Wage Data	Using Age Data	Using Wage Data	Using Age and Wage Data
Faculty	0.81	0.87	0.88	0.88	0.89
Graduate	0.09	0.73	0.73	0.73	0.73
Staff / Other	0.97	0.96	0.96	0.94	0.94
Postdoc	0.87	0.57	0.63	0.70	0.63
Undergrad	0.23	0.17	0.36	0.39	0.37
Overall	0.65	0.82	0.84	0.84	0.84
Fraction of job titles whose mode of predicted classes matches the true class by true class					
Faculty	0.81	0.90	0.90	0.92	0.92
Graduate	0.09	0.13	0.12	0.12	0.12

Staff / Other	0.97	0.95	0.97	0.95	0.96
Postdoc	0.87	0.85	0.87	0.86	0.86
Undergrad	0.23	0.10	0.08	0.10	0.09
Overall	0.65	0.78	0.79	0.79	0.79

E. Transitional and Concurrent Titles

As indicated, people can hold multiple titles at a point in time (or in close succession) and can transition between titles. As some transitions are more common than others (i.e., a transition from undergraduate to graduate to postgraduate to faculty is more common than the reverse set of transitions), it is possible to use transitions between titles and concurrent titles (more precisely, occupational classes that are associated with these titles) as predictors in the random forest classifier to improve predictive accuracy. Transitional and concurrent titles can also be used to identify unlikely transitions in the “ground-truth” data, providing an opportunity for a revision. Beyond improving the accuracy of the data, exploring concurrent positions and transitions can add to the richness of our data by providing information on career paths.

Using concurrent job titles and the transitions between job titles requires some form of iterative procedure. Obviously, the complete mapping between the set of job titles to itself is too high dimensional to be of any practical use. Thus, we use the following approach. In the first iteration, we predict occupational class using only job titles as predictors. In the second iteration, the predicted classes of the transitional/concurrent

titles from the first iteration are used as predictors, along with the job titles. In principle, this process could be iterated until the prediction converges according to some criterion. The example below, where an individual held three job titles in sequence, illustrates our approach.

Table A6. Illustration of iterative procedure using transitions

First Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	Prediction
Student Help			-	Staff
Research Assistant	-		-	Undergraduate
Postdoctoral Researcher	-			Postgraduate

Second Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	New Prediction
Student Help			Undergraduate	Undergraduate
Research Assistant	Staff		Postgraduate	Graduate
Postdoctoral Researcher	Undergraduate			Postgraduate

Third Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	New Prediction
Student Help			Graduate	Undergraduate
Research Assistant	Undergraduate		Postgraduate	Graduate
Postdoctoral Researcher	Graduate			Postgraduate

In this example, “postdoctoral researcher” is pivotal. Because the job title is so informative, its predicted class during the second iteration is not affected by the wrong prediction for the preceding title (i.e., it is unlikely to transition directly from undergraduate to postgraduate, but it is even more unlikely for a non-postgraduate student to have a job title “postdoctoral researcher”).

Of course, the time gap between the consecutive titles should also be taken into account. For the above example, if the time gap between “research assistant” and “postdoctoral researcher” is more than several years, the initial prediction of undergraduate for the job title “research assistant” may be more appropriate than the revised prediction of graduate. Here, we do not leverage the time gap between job titles in the model, but do include age, which contains somewhat similar information regarding the timing of job titles.

One issue with the iterated prediction procedure is that the convergence is not guaranteed, especially when there is no pivotal job title. For example, consider the following individual who held two job titles simultaneously.

Table A7 – Illustration of iterative procedure to use concurrent titles

First Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	Prediction
Tutor		-		Graduate
Grader		-		Undergraduate

Second Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	Prediction
Tutor		Undergraduate		Undergraduate
Grader		Graduate		Graduate

Third Iteration:

Job Title	Preceding Class	Concurrent Class	Succeeding Class	Prediction
Tutor		Graduate		Graduate
Grader		Undergraduate		Undergraduate

Because we cannot say whether a tutor or a grader is definitely an undergraduate or graduate, it is possible that, when making a revised prediction, the random forest classifier will simply adopt the classification for the concurrent title predicted in the previous iteration. As a result, the prediction will flip-flop and the algorithm will never stop. Of course, the presence of other people in these occupations mitigates this problem at least to some extent.

Data Construction

As a first step toward incorporating transitional and concurrent classes in the random forest classifiers, we included the manually classified transitional and concurrent classes in our training data in the model rather than predicted occupations. The resulting predictive accuracy is expected to provide an upper bound for the accuracy obtained from the iterated procedure described above.

To construct our sample, the monthly transaction records were collapsed at the individual-title-year level. That is, for each individual, for each calendar year, for each job title held during the year, we kept the individual-title-year record if the individual appeared in the transaction record both before July 1 and after September 30 with the job title. This is to avoid potential noise in the data when annual income is merged. Suppose an undergraduate student held a research assistant position from January through June. Then he or she graduated and obtained a full-time job. If the individual was included in our sample, it would appear that the annual income of the individual is too high to be an undergraduate, and it can potentially mislead the random forest classifier.

The concurrent job titles are defined to be a group of job titles that were held by an individual within a year. When there were multiple concurrent job titles, we selected the one for which the individual was paid the longest. The preceding job title is defined to be a job title held by an individual in the years preceding the current year. When there were multiple preceding job titles, we selected the most recent one. The succeeding job title is defined to be a job title held by an individual in the years succeeding the current year. When there were multiple succeeding job titles, we selected the one that immediately followed the current job title. Because we restricted our sample to individuals appearing in the transaction data between 2012 and 2014, the occurrence of multiple concurrent or transitional job titles was rare.

Before fitting the random forest classifier, transitional and concurrent classes were binarized because the random forest classifier cannot process categorical data. Each of preceding, concurrent, and succeeding class variables was decomposed into five indicator variables (faculty, postgraduate, graduate, undergraduate, and staff/other).

Methodology

To properly measure the effect of including transitional/concurrent classes on the predictive accuracy, we created the following subsets of observations:

- Everyone: Every observation
- None: Observations without any transitional or concurrent titles

- Prec: Observations with preceding title (may or may not have succeeding or concurrent titles)
- Succ: Observations with succeeding title (may or may not have preceding or concurrent titles)
- Conc: Observations with concurrent title (may or may not have preceding or succeeding titles)
- Any: Observations with at least one of preceding, succeeding, or concurrent title (can have multiple)

We expect that inclusion of transitional/concurrent classes have no effect on occupations where no observations have any transitional/concurrent classes while it will have the largest effect on occupations with many cases with all of the three classes. Each of the six subsets listed above served as a test set, and the random forest classifier was fitted with and without transitional/concurrent classes as predictors.

Regarding the training set, it is unclear whether the set should be restricted in the same way as the test set. Consider the test set “Prec”. On the one hand, it seems reasonable to restrict the training set to only observations with preceding titles. This is because if observations without preceding title were to be included in the training set, the importance of the preceding class in predicting the current class may be discounted. On the other hand, requiring observations in the training set to have preceding titles greatly reduces the number of qualified observations, possibly leading to overfitting. Since the

effect of restricting the training set is unclear, we fitted the random forest classifier with and without restriction on the training set.

As shown in Table A8, including the concurrent and transitional occupation has minimal effect on the predictive accuracy. This is most likely due to the limited number of relevant observations in the training set; therefore, as more universities participate in the IRIS project and provide data over a longer time period, the concurrent and transitional occupation may become a useful predictor.

Table A8

Features	Training Set	Everyone	None	Prec	Succ	Conc	Any
Job title	Unrestricted	0.83	0.83	0.82	0.84	0.83	0.83
Job title and occupation	Unrestricted	0.82	0.82	0.83	0.85	0.82	0.84
Job title	Restricted to relevant group	0.83	0.82	0.83	0.82	0.86	0.83
Job title and occupation	Restricted to relevant group	0.82	0.83	0.83	0.85	0.78	0.83