

Supplementary Materials

1 Supplementary data on faculty rank

We report below additional information regarding the rank breakdown of faculty in our syllabi dataset ($n = 169$, Supplementary Table 1) and in the IPEDS data ($n = 2826$, Supplementary Table 2). For details on faculty rank definitions, see Section 3.3.2. We note that the faculty who submitted more than one syllabus to our dataset ($n = 39$) did so while holding the same rank.

Faculty by rank in syllabi dataset ($n = 169$)

Time Period	Professor	Associate Professor	Assistant Professor	Instructor	Total
Pre-COVID-19	7 (18.42%)	13 (34.21%)	6 (15.79%)	12 (31.58%)	38 (22.49%)
COVID-19	31 (23.66%)	32 (24.43%)	13 (9.92%)	55 (41.98%)	131 (77.51%)
Total	38 (22.49%)	45 (26.63%)	19 (11.24%)	67 (39.64%)	169

Supplementary Table 1. Breakdown of faculty by rank in our syllabi dataset ($n = 169$). Raw numbers (and corresponding percentages) of faculty with the rank of Professor, Associate Professor, Assistant Professor, and Instructor are shown for two time periods: pre-COVID-19 and COVID-19. Pre-COVID-19 considers syllabi for all courses taught between 2016 and 2019, as well as C-term (January to March) 2020. COVID-19 counts all courses taught from D-term (March to May) 2020 forward.

Faculty by rank in IPEDS dataset ($n = 2826$)

Time Period	Professor	Associate Professor	Assistant Professor	Instructor	Total
Pre-COVID-19	423 (26.74%)	394 (24.91%)	275 (17.38%)	490 (30.97%)	1582 (55.98%)

COVID-19	345 (27.73%)	252 (20.26%)	206 (16.56%)	441(35.45%)	1244 (44.02%)
Total	768 (27.18%)	646 (22.86%)	481 (17.02%)	931 (32.94%)	2826

Supplementary Table 2. Breakdown of faculty by rank in the IPEDS dataset ($n = 2826$). Raw numbers (and corresponding percentages) of faculty with the rank of Professor, Associate Professor, Assistant Professor, and Instructor are shown for two time periods: pre-COVID-19 and COVID-19. Pre-COVID-19 considers syllabi for all courses taught between 2016 and 2019, as well as C-term (January to March) 2020. COVID-19 counts all courses taught from D-term (March to May) 2020 forward.

2 Details on classification tree algorithm

In Section 3.4 of the manuscript, we discuss using a predictive classification tree algorithm to identify the most influential demographic factors when predicting certain inclusivity outcomes. We produce our classification trees using the R statistical package **rpart** which provides advanced tools for precise analyses and reliable results (Therneau and Atkinson, 2022).

The goal of our classification tree algorithm is to build trees whose last leaf is a *pure node*, i.e., a node where the y-variable prediction can be carried out with 100% accuracy. The *weighted entropy*, E , of a node is defined as:

$$E = - \sum_i \rho(x_i) p(x_i) \log(p(x_i)), \quad (\text{SM1})$$

and indicates the value or utility of a certain outcome, x_i , when selecting a node variable. Here, $p(x_i)$ is the probability of outcome x_i , and $\rho(x_i)$ is the n of the node x_i (Kelbert et al., 2017). Note that since the entropy is weighted by the n , the algorithm in some cases prefers variables with greater sample sizes. While it is not always possible to produce a pure node with $E = 0$, the classification tree algorithm chooses as a node variable the one that minimizes its entropy. This is accomplished by computing the entropy associated with each possible variable selection and then picking the smallest

one so that the *information gain* at each split is maximized. The information gain for each variable selection is computed by subtracting the weighted entropy of each branch generated from the parent's own weighted entropy.

We impose additional requirements to guarantee reasonable algorithmic decisions despite our small dataset ($n = 169$). Specifically, when choosing a variable for a split, each resulting branch must have at least three syllabi. Additionally, no variable is selected for a split (and the last node is simply a leaf) if the tree outcome cannot be improved enough no matter which variable is chosen. The algorithm computes a complexity parameter, cp , for each possible variable choice and ignores all options for which cp is not greater than 0.01. Imposing the complexity parameter threshold reduces overfitting and produces smaller trees with more easily interpretable results.

2.1 Example computation of complexity parameter cp for Figure 8

The complexity parameter cp is one of the stopping parameters in the classification tree algorithm, i.e., the tree stops splitting if no variable available can create a split with $cp > 0.01$ where each branch has at least 3 syllabi. Its use penalizes trees with numerous splits and prevents overfitting by ensuring trees improve significantly with each split. The formula for the complexity parameter is:

$$cp = \frac{\Delta_{relative\ error}}{\Delta_{number\ of\ splits}}. \quad (SM2)$$

The *relative error* is the number of incorrect predictions performed by the classification tree when choosing a node variable divided by the number of incorrect predictions the tree would make if there was no split.

We compute the complexity parameter cp for our inclusivity statement classification tree reported in Figure 8. We have 57 syllabi with inclusivity statements and 112 without (for a total of

$n = 169$ syllabi). Without any splits the tree would predict “no” as the majority class (since $112 > 57$), resulting in 57 incorrect predictions. This leads to a relative error of $57/57 = 1$. First, the tree splits on gender lines, incorrectly predicting 30 syllabi on the “Men” side (left, $123 - 93 = 30$). On the “Women” side (right), the tree then isolates the Sciences (SCI) field, resulting in 1 incorrect prediction ($12 - 11 = 1$). Of the remaining 34 syllabi after the SCI split, 26 have inclusivity statements, 8 do not. Without any other splits, the tree would predict “yes” as the majority class (since $26 > 8$), resulting in 8 incorrect predictions. Hence, after 2 splits, we have a total number of incorrect predictions of $30 + 1 + 8 = 39$. Thus, we have $cp = (1 - 39/57)/2 \sim 0.1579 (> 0.01)$.

Finally, we consider the last split in the tree. Separating Engineering (ENG) from the remaining fields on the right side of the tree brings the number of incorrect predictions to $30 + 1 + 1 + 5 = 37$, and the change in splits is now 3 (splits) $- 2$ (splits) $= 1$. The complexity parameter for the last split in the tree is $cp = (39/57 - 37/57)/1 \sim 0.0351 (> 0.01)$.

3 Supplementary data on syllabi with no Identity Safety Cues

We report the number of syllabi in our dataset with no Identity Safety Cues ($n = 94$), grouped by faculty rank (Supplementary Table 3) and by faculty field (Supplementary Table 4). For details on faculty rank and field definitions, see Sections 3.3.2 and 3.3.3, respectively.

Syllabi in our dataset with no Identity Safety Cues by faculty rank ($n = 94$ of 169)

Professor, $n = 38$	Associate Professor, $n = 45$	Assistant Professor, $n = 19$	Instructor, $n = 67$	Total, $n = 169$
18 (47.4%)	31 (68.9%)	8 (42.1%)	37 (55.2%)	94 (55.6%)

Supplementary Table 3. Breakdown of syllabi in our dataset with no Identity Safety Cues (ISCs) ($n = 94$ of 169) by faculty rank. Raw numbers (and corresponding percentages) of syllabi with no ISCs

are shown grouped by authors’ rank: professor, associate professor, assistant professor, and instructor.

Syllabi in our dataset with no Identity Safety Cues by faculty field (*n* = 94 of 169)

ENG <i>n</i> = 24	HUA <i>n</i> = 35	MATH <i>n</i> = 38	SCI <i>n</i> = 32	SOS <i>n</i> = 22	TECH <i>n</i> = 18	Total <i>n</i> = 169
12 (50.0%)	8 (22.9%)	23 (60.5%)	24 (66.7%)	12 (54.5%)	15 (83.3%)	94 (55.6%)

Supplementary Table 4. Breakdown of syllabi in our dataset with no Identity Safety Cues (ISCs) (*n* = 94 of 169) by faculty field. Raw numbers (and corresponding percentages) of syllabi with no ISCs are shown grouped by authors’ field: Engineering (ENG), Humanities & Arts (HUA), Mathematical Sciences (MATH), Sciences (SCI), Social Sciences (SOS), and Technology (TECH).