

Bias Audit for Plum's Talent Match Model

plum



Table of contents

Table of contents	2
Bias evaluation for Plum’s Talent Match Model	2
Executive summary	2
Plum Talent Match Model system overview	3
System description	3
Bias evaluation results	3
Data used	3
Scoring methodology	4
Univariate Categories	4
Results by gender	4
Results by race	5
Intersectional Categories	5
Results by race and gender	5
Appendix	6
Report applicability	6
About FairNow	6

Bias evaluation for Plum’s Talent Match Model

Report prepared by FairNow on October 28, 2025

Executive summary

The Plum Talent Match Model uses assessments of cognitive and personality traits to match job seekers with jobs. On the platform, job seekers answer a series of questions that measure various traits, and employers determine the most relevant traits associated with each of their job postings. The Plum Talent Match model outputs a score that measures the alignment between the job seeker’s assessment and the associated traits of a job requisition.

Testing in this report focuses on the Plum Talent Match Model’s ability to score the alignment between a job seeker and a posted job based on cognitive and personality traits. Performance was assessed for gender, race, and gender–race combinations.

Test results based on the scoring methodology outlined in this report do not find evidence of adverse impact, as all groups had selection rates at least 80% of the most favored group in accordance with the EEOC's four-fifths rule.

Plum Talent Match Model system overview

System description

The Plum Talent Match model is designed to enable organizations to determine how well suited a candidate is for a role by comparing a candidate's set of Talents to those required for success in the role. Talents are derived from the Five-Factor Model (FFM) of personality as well as facets of cognitive ability, and are grounded in years of research on personality, cognitive ability, and social intelligence. The following inputs are used to identify Talents for matching purposes:

- A set of expert contributors within the employer's organization completes Plum's Match Criteria Survey to define behavioral role requirements. Outcomes of the survey are aggregated to determine the top 5 Talents most critical for success in the role.
- Candidates are asked to complete the Plum Discovery Survey, which includes questions focused on personality, problem solving, cognitive ability and social intelligence. From this survey, the candidate receives scores for each of the 12 Talents.

The Talent Match model combines candidate Talent results and the ranked Talents from the employer to calculate scores reflecting the job seeker's alignment with the employer's top criteria for a given role. Candidates are ranked based on these scores, which can be used to influence whether an employer decides to move a job seeker forward in the hiring process.

Bias evaluation results

Data used

This evaluation includes historical data from applicants assessed by Plum's Talent Match model from September 14, 2024 to August 29, 2025. This date range was chosen to include all samples after the cutoff date of the data used in Plum's last bias audit. Demographic data for this population was collected by Plum. Candidates had the option to submit demographic data directly to Plum when they began an assessment on the platform. Results were tested only for candidates where demographic data was available.

Gender was represented as one of the following values:

- Female

- Male
- Non-binary
- Opt-out or otherwise unknown (not included in the analysis)

Race/Ethnicity was represented as one of the following values:

- Asian
- Black
- Indigenous
- Latinx
- Multiple Races
- White
- Opt-out or otherwise unknown (not included in the analysis)

Scoring methodology

As part of the candidate assessment process, each application receives a score ranging from 30 to 99 based on how well the candidate's Talents align with the top 5 Talents specified by the hiring organization. All candidates are shown to the recruiter, ranked by match score. Because there is not a designated pass/fail cutoff, this audit applies a scoring rate method recommended by NYC Local Law 144 and compares scores against a median value to determine impact ratios.

NYC Local Law specifies that a median value should be calculated across the "full sample of applicants". Because the Plum Talent Match model produces scores that vary widely at the job level, FairNow calculated a median value across the full sample of applicants for each job, i.e., evaluating candidate outcomes relative to others applying for the same job versus all candidates overall. It is FairNow's perspective that this approach most accurately reflects how Plum's Talent Match Model is used in practice.

In supplemental testing using a single median value across all jobs, FairNow found impact ratios above 80% for all groups except Black (79%) and Black Female (77%). Differences between single-median and job-level median outcomes may arise from composition effects: different jobs can have distinct baseline scores and demographic compositions, which can produce aggregate-level differences that are not seen within individual jobs.

Univariate Categories

Results by gender

Gender	# of Applications	# Selected	Scoring Rate	Impact Ratio
Female	26,477	13,844	52%	97%
Male	52,560	28,268	54%	100%
Non-binary	151	74	49%	N/A

“N/A” refers to categories that comprised less than 2% of the total. These groups were excluded from the analysis.

There were 477,877 applications for which the candidate’s gender was not known. This data was not included in the above table.

Results by race

Race	# of Applications	# Selected	Scoring Rate	Impact Ratio
Asian	43,293	23,257	54%	98%
Black	8,861	4,689	53%	97%
Indigenous	123	72	59%	N/A
Latinx	5,051	2,676	53%	97%
Multiple Races	1,927	1,051	55%	100%
White	11,065	5,669	51%	94%

“N/A” refers to categories that comprised less than 2% of the total. These groups were excluded from the analysis.

There were 486,745 applications for which the applicant’s race/ethnicity was not known. This data was not included in the above table.

Intersectional Categories

Results by race and gender

Race	Gender	# of Applications	# Selected	Scoring Rate	Impact Ratio
Asian	Female	14,261	7,568	53%	98%
Asian	Male	28,844	15,583	54%	100%
Asian	Non-binary	56	32	57%	N/A
Black	Female	3,001	1,520	51%	94%
Black	Male	5,842	3,160	54%	100%
Black	Non-binary	4	2	50%	N/A
Indigenous	Female	42	22	52%	N/A
Indigenous	Male	74	44	59%	N/A
Indigenous	Non-binary	6	5	83%	N/A
Latinx	Female	1,610	854	53%	98%
Latinx	Male	3,404	1,795	53%	97%
Latinx	Non-binary	8	3	38%	N/A
Multiple Races	Female	587	313	53%	N/A
Multiple Races	Male	1,281	704	55%	N/A
Multiple Races	Non-binary	15	11	73%	N/A
White	Female	3,361	1,680	50%	92%
White	Male	7,616	3,959	52%	96%
White	Non-binary	54	15	28%	N/A

“N/A” refers to categories that comprised less than 2% of the total. These groups were excluded from the analysis.

There were 486,991 applications for which the applicant’s gender or race/ethnicity were not known. This data was not included in the above table.

Appendix

Report applicability

This evaluation tests for adverse impact in accordance with the EEOC's four-fifths rule. It represents an independent bias audit in alignment with methodologies from New York City Local Law 144 and constitutes algorithmic discrimination testing for race, gender, and race–gender combinations that may be relevant for the Colorado AI Act, the European Union AI Act, and the California Civil Rights Council's employment regulations on AI.

Results in this report represent a point-in-time snapshot and are based on the relevance of the testing data available at the time of generation. We encourage ongoing monitoring over time.

About FairNow

FairNow is an organization dedicated to helping companies leverage AI in a responsible, fair, and well-managed way. FairNow is an independent auditor in alignment with the specifications of New York City Local Law 144.