



Greenhouse Talent Matching NYC LL 144 Audit Report

Generated from

Greenhouse AI Assurance Dashboard

July 30, 2026

Table of Contents

1. REPORT SUMMARY

1.1 Audit information

2. ABOUT WARDEN AI

2.1 Company summary

2.2 Company information

2.3 Independence statement

3. AUDIT SCOPE

3.1 System details

3.2 Data details

3.3 Audit details

4. AUDIT RESULTS

4.1 Sex

4.2 Race/Ethnicity

4.3 Intersectional (Sex x Race/Ethnicity)

5. METHODOLOGY

5.1 Methodology overview

5.2 Disparate impact analysis

6. DISCLAIMER

Report summary

Warden AI is engaged by Greenhouse to perform independent, ongoing bias audits of the Talent Matching system. This report summarizes the most recent audit, generated using Warden AI's assurance platform and reviewed by our audit team.

The audit aligns with the requirements of the NYC Local Law 144, which mandates a bias audit of automated employment decision tools (AEDT). The methods used meet the specific requirements for conducting such an audit as published in the final rules of the NYC Department of Consumer and Worker Protection (DCWP). Warden's approach supports these aims while promoting fairness and accountability in automated decision-making.

The system was evaluated using a purpose-built test dataset, as sufficient historical data was not available. One bias-detection technique was applied to evaluate the system's behavior on this dataset at the time of testing.

This report reflects the system's behavior under controlled test conditions at the time of evaluation and should not be interpreted as a formal certification or compliance determination. The findings do not predict outcomes for any specific employer, job opportunity, or real-world deployment.

Warden's engagement does not include determining whether the audited system is legally classified as an AEDT or other regulated technology. Regulatory applicability depends on the system's particular configuration, deployment and use. The existence of this audit should not be interpreted as an admission or determination that any particular law applies.

Coverage of NYC LL 144

While NYC Local Law 144 sets out additional requirements such as notice to candidates and publication of audit summaries, this document focuses on the bias audit requirements for automated employment decision tools (AEDT). Accordingly, this report should not be interpreted as evidence of full compliance with NYC Local Law 144.

Report summary

Audit information

System audited:	Talent Matching
Audit frequency:	Monthly
Latest audit date:	July 30, 2026
Sample size:	16,214



About Warden AI

Company summary

At Warden AI, our mission is to reduce societal discrimination through fair and transparent AI. We provide third-party oversight into AI systems, building trust and increasing adoption.

We are an independent AI auditor and assurance platform that performs ongoing audits to ensure AI systems are fair, explainable, and transparent. Our team brings extensive experience across AI, regulation, and research, including industry and academia, to deliver our solution.

Our system integrates with the AI system that is under test, allowing for continuous testing and monitoring. Our methodology employs a combination of bias detection techniques and uses our proprietary datasets and/or historical data from the system.

Independence statement

Warden Technologies, Inc. is an independent AI audit and assurance provider. Fees associated with our service are solely for our evaluation and their payment is not related to the outcome of the results.

Our services are strictly limited to testing and monitoring the trustworthiness of AI systems. We do not form part of the solution or in any way affect how the system under test works.

Company information

Registered address:

600 Congress Ave, 14th Floor,
Austin, TX 78701, United States

Website:

<https://warden-ai.com>

Contact:

contact@warden-ai.com

Audit scope

AUDIT DATE

July 30, 2026

AUDIT FREQUENCY

Monthly

DATA SOURCE

Warden dataset

SAMPLE SIZE

16,214

TEST INPUTS

Job criteria, Candidate profile

TEST OUTPUTS

Match label

BIAS DETECTION TECHNIQUES

1

PROTECTED CLASSES

3

System details

Greenhouse's Talent Matching feature helps organizations to improve their hiring process by identifying the most qualified candidates based on specific criteria and powered by AI.

Candidate profiles are assessed against a set of customer-defined and weighted job criteria, such as their skills and experience. Each candidate is assessed against the same set of criteria individually, and an overall match score is produced to indicate how well the candidate fits the overall job criteria.

Data details

This audit used Warden's purpose-built test dataset to evaluate system behavior in a controlled setting. Historical candidate data with sufficient demographic coverage was not available.

The dataset contains real resumes collected through multiple sources, specifically curated for bias and fairness evaluation. All resumes are used with an appropriate legal basis, such as explicit participant consent. Demographic information such as the profile's sex is self-reported and standard data quality checks are applied before use.

The dataset also includes job criteria covering a wide range of occupations. These are provided either as full job descriptions or as structured criteria, such as skill requirements or evaluation rubrics. Resumes and job criteria are matched at the occupational level to avoid assessing unrelated roles.

Audit scope

Audit details

Audits are performed on a regular cadence to continuously monitor potential bias risks. The most recent audit was completed on *July 30, 2026* and evaluates potential bias across the following 3 protected classes:

- Sex, Race/Ethnicity, and Intersectional (Sex X Race/Ethnicity)

The following bias detection technique was applied:

- **Disparate impact analysis:** This analysis evaluates if the system is adversely impacting a protected group compared to another protected group. An impact ratio of 0.8 or higher is considered acceptable.

For the disparate impact analysis, the selection rate method was applied. Matches labeled as *Good Match* or *Strong Match* were classified as selected outcomes, while matches labeled as *Limited Match* or *Partial Match* were classified as not selected.

Audit results

Sex bias

Group bias (Disparate impact analysis)

Result: Clear

Sample size: 16,214

Group	Samples	Selected	Selection rate	Impact ratio
Female	8,610	3,100	36.00%	1.00
Male	7,604	2,722	35.80%	0.99

The results indicate equitable outputs across all groups.

Race/Ethnicity bias

Group bias (Disparate impact analysis)

Result: Clear

Sample size: 16,214

Group	Samples	Selected	Selection rate	Impact ratio
Asian	3,708	1,257	33.90%	0.89
Black or African American	3,129	1,187	37.94%	1.00
Hispanic or Latino	2,005	696	34.71%	0.91
White	7,372	2,682	36.38%	0.96

The results indicate equitable outputs across all groups.

Audit results

Intersectional (Sex X Race/Ethnicity) bias

Group bias (Disparate impact analysis)

Result: Clear

Sample size: 16,214

Group	Samples	Selected	Selection rate	Impact ratio
Asian / Female	1,816	628	34.58%	0.87
Asian / Male	1,892	629	33.25%	0.84
Black / Female	1,788	656	36.69%	0.93
Black / Male	1,341	531	39.60%	1.00
Hispanic / Female	1,156	394	34.08%	0.86
Hispanic / Male	849	302	35.57%	0.90
White / Female	3,850	1,422	36.94%	0.93
White / Male	3,522	1,260	35.78%	0.90

The results indicate equitable outputs across all groups.

Methodology

Methodology overview

Our methodology for evaluating AI systems is designed to ensure fairness and transparency. Our comprehensive approach includes ongoing auditing, multiple bias detection techniques, the use of diverse datasets, and human oversight.

This rigorous approach enables us to accurately report on the level of bias in the system and build trust with the system's users and stakeholders.

Black box testing

We use black-box testing techniques to evaluate AI systems. This approach examines the system's outputs in response to specific inputs without needing to understand the internal workings.

This enables us to make systematic judgements across different AI systems with different underlying models.

Ongoing audits

AI systems change frequently (often monthly, weekly, or even daily). Our audits are performed on a regular basis at the frequency detailed in this report. The exact frequency is determined with the AI provider based on the nature of their system and their propensity for product updates.

In addition to the scheduled evaluations, the AI provider can also choose to have an audit performed on-demand between scheduled audits if they have a significant product update.

Adherence to NYC Local Law 144

While our full bias auditing approach goes beyond the requirements, it is in adherence with NYC Local Law 144 of 2022, and meets the specific requirements for conducting a bias audit of automated employment decision tools (AEDT) as published in the final rules of the NYC Department of Consumer and Worker Protection (DCWP).

Our Disparate Impact Analysis identifies adverse impact on protected groups separated by sex and race/ethnicity as mandated by the Local Law 144.

Methodology

Hybrid auditing

Our evaluation process combines automated methods with human oversight to ensure accuracy and reliability.

By integrating AI systems with our standardized datasets, we can conduct large-scale and frequent audits. This approach is complemented by human-led data curation and quality assurance processes for creating the datasets. Additionally, our team of experts reviews and validates the results of audits to ensure reliability.

Diverse datasets

Our auditing framework uses a mixture of data. We have our own proprietary datasets which provide an independent benchmark of the AI system. Our dataset is formed of real data sourced from real people where consent has been provided.

Where applicable, we also use both historical and live data to provide context for the system's long-term performance and its current real-time operations.

All datasets are ethically sourced and we adhere to high standards of data collection practices. We are committed to maintaining confidentiality and protecting personal data. Some of our evaluations require datasets that contain elements of personal information to test specific AI functionalities. In such instances, we ensure that consent has been explicitly obtained for the use of this information.

Methodology

Disparate impact analysis

Disparate Impact Analysis evaluates whether a protected demographic group is adversely affected compared to other groups. This is achieved by comparing its selection rate to the most selected group. The goal is to ensure that the AI system does not disproportionately disadvantage any specific group based on inherent characteristics such as race, ethnicity, or sex.

Selection rate

Selection rate is a measure used to evaluate the proportion of individuals in a specific group who receive favorable outcomes from the AI system.

To calculate a group's selection rate, we divided the number of individuals who are considered selected by the total number of individuals with the group.

$$\text{Selection rate} = \frac{\text{Number of individuals within group who are considered selected}}{\text{Total number of individuals within group}}$$

Impact ratio

The impact ratio is a metric used to measure potential adverse impact on a group by comparing its selection rate to the most selected group.

$$\text{Impact ratio} = \frac{\text{Selection rate for the group}}{\text{Selection rate of the most selected group}}$$

An impact ratio of 1 indicates no adverse impact, whereas a lower ratio indicates a higher likelihood of adverse impact. According to the four-fifths rule, an impact ratio of 0.8 (80%) or higher is considered acceptable, indicating that the AI system's outcomes are equitable across different demographic groups.

Disclaimer

This report has been prepared by Warden Technologies, Inc. to provide an independent audit of the AI system developed by the AI provider in question, based on our proprietary methodologies and datasets. The results and conclusions presented in this report reflect our best judgments derived from the information available at the time of evaluation. While we strive for accuracy and completeness, we cannot guarantee that our evaluation is exhaustive or that there are no errors.

Our methodology is designed to identify potential issues of bias and other trust factors in the AI system under examination. However, our approach, like any evaluation methodology, has its limitations. It is important to understand that our findings do not guarantee the absence of any bias, flaws, or limitations within the audited AI system. Instead, they indicate that, based on our specific testing framework and within the scope of our analysis, no significant issues were identified.

This report is intended for informational purposes only and should not be interpreted as a guarantee of the system's performance, fairness, or suitability for any specific purpose or use case. Warden Technologies, Inc. disclaims any liability for any decisions made or actions taken based on the information provided in this report. By using this report, the reader agrees to assume all risks associated with such decisions or actions.



Greenhouse Talent Matching NYC LL 144 Audit Report