

AI-Driven Recruitment Readiness Among Indian Graduates: Evidence on Skills, Training, Perceived Difficulty, Confidence, and Screening Outcomes

Chitrika Nare¹, Ravindra Gharpure², Rahul Mohare³

¹Research Scholar, CIBMRD, RTM Nagpur University

²Assistant Professor, CIBMRD, RTM Nagpur University

³Assistant Professor, Ramdeobaba University, Nagpur

Abstract

AI-facilitated recruiting is beginning to seep into graduate hiring pipelines but most research to date focus either on employers' adaptation to AI recruitment processes or on algorithmic fairness or generally on digital skills rather than on applicant preparedness to it. In this cross-sectional study, we test whether AI-related skill set and/or specific formal training (H1-2), city tier (H3), type of educational institution (H4), past recruitment exposure (H5), perceived barrier factors (H6-8) are associated with perceived recruiting difficulty (PD), confidence (CC), and screen in clearance (ClearC). Using descriptive comparative statistics with effect measures (e.g. Medians; clearance percentages) as well as Mann Whitney U test; Spearman correlation; chi square tests along with Cramer's V ; Kruskal Wallis test and Cliff's delta test. An eleventh hypothesis that institution type proxy was significantly associated with CC did attain to nominal significance ($p=0.037$) but this did not pass through a stringent Bonferroni adjusted threshold (0.006) and was estimated using an indirect proxy which proxies a factor which could be identified direct in research design. Our findings reveal a surprising lack of correlation between standard preparedness factors and how confidently applicants approach their future careers; or whether they struggle in the recruitment process or pass the screening process. Such significant non-findings will be interpreted in terms of potential limitations. This may add, tentatively, some applicant-centric empirical evidence into research literature regarding preparedness to AI-recruiting in Indian graduate setting, pointing towards avenue for subsequent studies e.g. By using stronger methods as multiple regression and ordinal regression analysis and by directly specifying of type of institution for H4.

Keywords: AI-driven recruitment; graduate employability; technology readiness; digital divide; algorithmic screening; India; perceived recruitment difficulty; recruitment confidence

1. Introduction

AI is gradually becoming a part of the recruitment process at almost all stages resume screening automated filtering online asynchronous interviews and the machine-driven selection of candidates. The change is not only in the method of hiring but is also a change in the knowledge and preparedness requirements of the job seekers. Whereas the past research on employability focused on curriculum-related skills, internships, and broad graduate competencies, a new dimension of preparedness that is less studied has emerged and comprises the applicants' ability to negotiate the AI-mediated screening: to what extent they know what goes on, if they feel confident in confronting it, and if they have the particular set of competences that such systems can value.

Previous studies on AI in recruitment mostly dealt with the employer and AI systems sides of the relationship such areas as algorithmic bias and the ways to combat it (Raghavan et al., 2020), fairness perceptions of candidates facing AI interviews (Acikgoz et al., 2020), and the effect of synchronised, machine-controlled, online video interviews on applicant behaviour (Suen et al., 2019). The other side of the problem, an individual's readiness for recruitment in AI-assisted environments (i.e. skills, training, past exposure, perceived obstacles that are likely to

affect an applicant's experience of AI-mediated recruitment and self-efficacy with regard to being successful in it), has received only scarce attention so far.

Such a mismatch may be most damaging for developing countries like India where higher education turns out a massive, diverse group of graduates who are scattered in cities with varying degrees of technological infrastructure, resource endowment of educational institutions, and labour market density. The digital divide (van Dijk, 2006) indicates that inequality in the outcomes of technology usage is not simply a matter of having access to devices but covers the entire spectrum of motivations, skills, and actual use which is a very relevant way of thinking about how graduates from Tier-1, Tier-2, and Tier-3 cities may be able to deal with recruiting processes mediated through AI. Technology Readiness theory (Parasuraman, 2000) suggests further that besides objective access to technologies, people's readiness is influenced also by their predisposition of their personality toward a new technology. That means how willing (personality) people are for technology will define their interaction or usage of technology with a higher or lower level of difficulty.

Employment Theory (Yorke, 2006) advises against equating job-readiness with mere knowledge/skill ownership, rather, it defines job-readiness as a multi-perspective and context-dependent phenomenon.

This paper exploits a dataset that was constructed primarily for the study of various factors related to the recruitments such as candidates' exposure, their perception of the levels of difficulty associated with the interview/skills assessment, the skills they bring to the employers' table, the availability of training that could help them prepare, obstacles to getting hired, the employers' screening decisions, and graduates' perceptions of whether they have the needed confidence to meet job requirements.

We present the results of testing eight hypotheses in this paper covering these constructs. The issue of graduate readiness in a recruitment process mediated by artificial intelligence or a different form of AI, remains an unsettled and open question requiring further policy attention rather than a solution. As a consequence, this paper does not merely assume AI-recruitment readiness is a function of skills and training and considers how skills, training, motivation, barriers and job market factors influence ready graduates for an AI-recruiting process, it also reports the findings even when there are no significant results obtained.

1.1 Research gap, problem, and contribution

Based on the aforementioned shortcomings, the research problem of this paper is: to what degree are Indian students ready to understand AI-mediated recruitment and what is the relationship, if any, between confidence, perceived recruitment difficulty, and screening outcomes on various aspects such as AI-related skills, formal training, City Tier, institutional setting, recruitment experience, and perceived barriers? This paper's contribution is threefold: first it brings in applicant-side data on a large, graduate labour market that still uses AI technology for their recruitment; second it allows for all the variables to be discussed: skills, training, exposure, barriers, confidence, difficulty, and outcomes together in a single analytical framework that could have previously been isolated; and third it pinpoints, using full methodological disclosure, which of the commonly assumed readiness factors that were found to have only limited explanatory power in the sampled data will serve as evidence for the further design of more powerful predictive and multivariate studies). Rest of the paper is arranged according to the following sections. Section 2 provides thematic literature review.

2. Literature Review

2.1 AI-driven recruitment

Companies are slowly integrating AI tools such as resume screening algorithms, AI chatbots, and automated video candidate assessment tools, into their hiring processes mainly with the aim of reducing human error.

Raghavan, Barocas, Kleinberg, and Levy (2020) revealed how sellers of pre-employment evaluation tools communicate their fairness-ensuring arguments and practices, pointing out that industry methods for fairness are diverse. Their standards for fairness are often vague in comparison to the standards of the law on disparate-impact.

The main focus of these studies is to see AI recruitment as a design and governance issue. That is, these studies show that how a system is created and by whom is a main concern. They leave open how, practically, applicants deal with and make themselves ready for these systems.

2.2 AI and graduate employability

Graduate employability has never really been seen as a single result or the product of a degree curriculum (Yorke, 2006). In a hiring environment mediated by AI, this means that employability nowadays most likely includes not only disciplinary and generic skills but also one's own comfort with being evaluated algorithmically. Yorke's analysis is a warning against the idea that only one kind of input be it even AI skills will always lead to better job results. This caution is echoed in the study reported later in this paper where a largely null pattern of association is shown.

2.3 AI skills and recruitment readiness

A reasonable thought in policy and curricular debate is that visible AI-related skills (like being at ease with data tools, or grasping simple machine-learning concepts) should make the applicant see AI-driven recruitment as less formidable and at the same time raise their chances of passing algorithmic screening. Although the logic underpinning this is compelling, the claim itself has not yet been fully supported through a comparison with the difficulties and clearance outcomes reported by the applicants of Indian graduate programmes as the present study, via H1 and H4, seeks to do.

2.4 Training and certification

Formal training and certification are usually seen as ways to help get a job by improving skills or knowledge in a particular field. These ideas fit well with the Technology Acceptance Model of Davis (1989) generally speaking, it's the model that predicts perceived ease of use which, ultimately, results in the feeling of comfort when dealing with technology. Therefore, it might be argued that the more a person is prepared, the more their perceived ease of use and therefore the comfort level will increase while dealing with a technology mediated task.

On the other hand, whether getting more formal training and certification in general is really what boosts a job applicant's confidence to go through a hiring process managed by an AI system, is a question that the H2 research hypothesis seeks to answer with a real life test case study.

2.5 Digital inequality and city/institutional context

Dijkvan (2006)'s digital divide framework separates the motivational, the physical, the skills, and usage access, and cautions against a simplistic view of digital inequality as a black and white situation. Recruitment-wise, that would point to city-tier or institutional differences in readiness for AI-recruitment, as a phenomenon if they do exist, would work through several different access dimensions that could cause problems rather than just one resource shortage.

This research explores the link between city-tier and difficulty in understanding (hypothesis 3) and the relationship between institution-type proxy and confidence and difficulty (hypothesis 7). While doing so, they remain wary about taking it for granted that social inequality of the environment will show up in these exact types of result measures.

2.6 Recruitment exposure

It is also frequently supposed that being familiar with a selection process through repeated exposure leads to greater confidence. This is very much the case too for technology-readiness theory by Parasuraman (2000) which holds that people who interactively take part in using technology become more comfortable with it over time. H6 directly tests whether previous exposure to AI recruitment in particular is a significant predictor of higher level of confidence or lower perceived difficulty among Indian graduates.

2.7 Barriers and applicant confidence

There are several barriers a candidate could run into during AI-mediated hiring. These are: a lack of help and instructions on the use of the technology, technology inequality, and the fact that it is not fully explained or clear how the decisions by the algorithm were made. From a study that examined applicants' perceptions of AI-based selection (Acikgoz, Davison, Compagnone, & Laske, 2020), we get that in general, the AI-based interviews were seen as being less procedurally and interactionally fair than the human ones. One of the reasons for this was that AI-based interviews allowed candidates far less chance of explaining or showing their abilities.

This piece of theory inspired us to expect, via H5 and H8, that barrier categories and being technologically disadvantaged will correspond to decreased self-confidence as perceived by applicants.

2.8 Screening and AI-mediated selection

Studies have explored how AI-based video interviews operate for example, Suen, Chen, and Lu, (2019) looked at how AI-based synchrony analysis presence influences interview ratings and applicants' feelings, stressing the point that AI screening does a lot more than a neutral backdrop: it changes both results & opinions. Another experimental research in highly automated job interviews revealed that people's tolerance to automated processes is not constant; it changes depending on the stakes (Langer, König, & Papathanasiou, 2019). High-stakes contexts lead applicants to be more cautious than in training contexts. In a number of experiments involving algorithm-based recruitment, job applicants were repeatedly of the opinion that automated selection was not fair they ranked algorithm-only selection lower than human or human-assisted AI processes (Lavanchy, Reichert, Narayanan, & Savani, 2023).

Regardless of the outcome, either positive or negative towards the applicant, this tendency still held. This work is H4, an extension of the literature on this topic.

H4 asks the question of whether self-rated AI expertise is actually a predictor of getting selected as the right candidate in the current sample, i.e., it connects a capability measurement with an actual output measurement rather than just a mental outcome measurement.

3. Theoretical Framework

For three reasons, this work has been drawn from three different theoretical angles: first, these concepts have been measured; second, they are being included for reasons other than just breadth alone.

Technology Readiness theory (Parasuraman, 2000) characterises a person's tendency to adopt new technologies in terms of optimism, innovativeness, discomfort, and insecurity. It provides a reason, or motivation, for a broader readiness disposition rather than merely being a result of possessing skills such as confidence and perceived difficulty.

The concepts underlying the Digitale Divide are such that, if a difference were taken in these four factors: one's motivation, one's physical ability, one's level of skills, and, finally, usage access, it would lead to a situation wherein one can explain an inequality of access. The city-tier and technology-access barrier hypotheses (H3, H8) are motivated by the Digitale Divide framework (van Dijk, 2006).

Yorke (2006) proposes employability to be a concept which graduates' employability could be interpreted as a context-dependent achievement which has several dimensions. This theory also warns that if you are assuming that there is only one predictor of employability skills, training, or exposure and that this one predictor would directly result into your getting better recruitment results, then you need to reconsider.

Apart from that, the fact that such an assumption would be contrary to our data findings is, we think, also quite supportive to the Employability theory. This is because such an assumption is not in line with the data presented in chapter 8.

Besides, the Technology Acceptance Model (Davis, 1989) has been employed as an additional tool primarily due to its connection between perceived ease of use and the central perceived-difficulty construct shared by H1, H3,

and H6. The literatures on reaction to applicants and algorithmic hiring (Acikgoz et al., 2020; Langer, König, & Papathanasiou, 2019; Lavanchy et al., 2023; Raghavan et al., 2020; Suen et al., 2019) are brought in to highlight the research gap from the applicant's side and to make sense of the barrier and screening-findings issues instead of providing additional directly tested hypotheses. Basically, the theoretical propositions are regarded only as inspiring elements (not confirming) and therefore as theoretical reasons that the relationship under study is plausible whereas the statistical test in Section 8 is what actually tells us whether the data support the relationship.

4. Conceptual Framework

The conceptual model (see Figure 1) is a four-cluster structure that categorizes the study's predictors, including contextual factors (city tier, proxy of institution type), capability factors (skills related to AI, AI proficiency, formal training), experience factors (AI-recruitment exposure, applications per person), and barrier factors (guidance/training, technology access, system transparency), as well as associatively linking them with three outcome constructs: perception of difficulty in recruitment, confidence in AI-screened recruitment, and screening clearance. The arrows in the diagram show the hypothesized pathways that were tested bivariately, not confirmed causal links; therefore, the model must be interpreted as a representation of the eight hypotheses (H1, H8) tested in the study, not as a structural or path model.

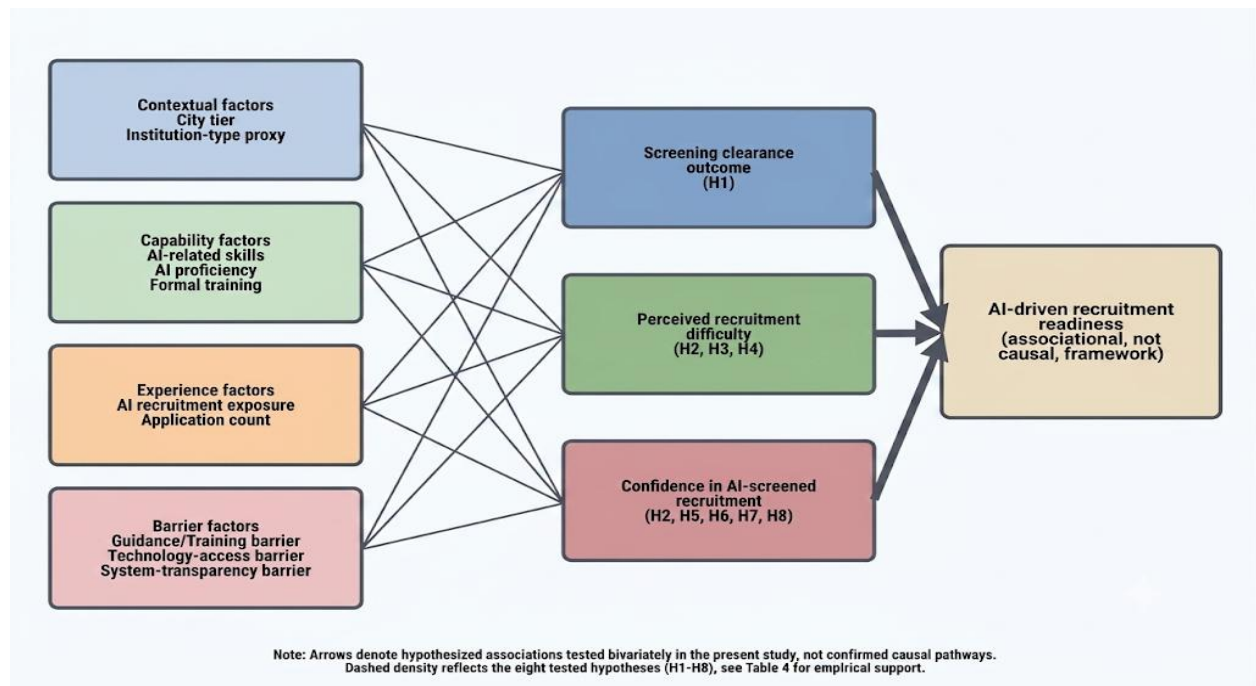


Figure 1. Conceptual framework of AI-driven recruitment readiness among Indian graduates.

5. Research Questions

- RQ1. What level of perceived difficulty do Indian graduates report regarding AI-driven recruitment?
- RQ2. Which barriers are most frequently reported during AI-based hiring?
- RQ3. Are AI-related skills and proficiency associated with perceived recruitment difficulty and screening clearance?
- RQ4. Is formal training associated with confidence in AI-screened recruitment?
- RQ5. Does contextual background — city tier or institution-type proxy — relate to recruitment readiness?
- RQ6. Are recruitment exposure and perceived barriers associated with confidence and perceived difficulty?

6. Hypotheses

Eight directional hypotheses (H1, H8) were generated from the previous research questions and they were assessed by the available set of data. All of them are mentioned, checked, and determined in Section 8. The list of hypotheses can be found in Part B but they will not be restated here in order to refrain from repetition; all statistical data including test statistics, p-values, effect sizes, and decisions for each hypothesis are provided in Table 4.

7. Methodology

7.1 Research design

This is a quantitative, cross-sectional, associational study. Data collection happened through one time point by Indian graduate respondents concerning their exposure level and opinions on AI-powered recruitment. Even though the study setup allows uncovering of statistically significant associations among variables it however doesn't allow making causal inferences.

7.2 Population and sample

The study aims at Indian graduate candidates either who have been directly involved in, or at least familiar with recruitment processes in today's context.

We did not fix the number of respondents or response rate to be achieved through the online questionnaire. A random sampling frame consisting of different categories such as male, female, different age groups, various graduation streams with different graduation years etc was chosen for drawing samples from.

The basic group-level statistics (average, median and percentage) as well for different subsections were available. See Section 8 of our report for more detailed analysis of our samples.

7.3 Sampling

The sampling procedure is a simple random sampling

7.4 Variables

Table 2 summarizes the variables analyzed, their role in the study, their measurement, and their level of measurement.

Variable	Role	Measurement / coding	Level
Perceived difficulty	Outcome	Single item, 1 = easy to 5 = difficult	Ordinal
Confidence	Outcome	Single item, 1 = low to 5 = high	Ordinal
Screening clearance	Outcome	Cleared vs. not cleared screening round	Binary categorical
AI-related skills	Predictor	Self-reported possession of demonstrable AI-related skills (yes/no)	Binary categorical
AI proficiency	Predictor	Self-rated proficiency: novice / intermediate / advanced	Ordinal categorical
Formal training/certification	Predictor	Completion of formal AI-related training or certification (yes/no)	Binary categorical
City tier	Predictor / context	Tier-1, Tier-2, Tier-3 (India classification, source not further specified)	Nominal categorical
Institution-type proxy	Predictor / context	Engineering/technology proxy vs. other-institution proxy, inferred from college name text	Binary categorical (proxy)

AI recruitment exposure	Predictor	Prior encounter with AI-driven recruitment (yes/no)	Binary categorical
Application count	Predictor	Number of job applications submitted (respondent-reported count)	Continuous / count
Primary barrier category	Predictor	Respondent's single selected primary barrier (guidance/training, technology access, system transparency, other)	Nominal categorical

Table 2. Variable definitions and coding.

7.5 Measurement

Perceived difficulty and confidence were measured in single-item 1-to-5 scales (1 = easy/low to 5 = difficult/high) matching the axis labelling on the original descriptive visualizations.

AI-related skills, formal training, and AI-recruitment exposure were variables coded as yes/no. AI proficiency was an ordinal category variable (novice, intermediate, advanced).

City tier was a three-level nominal variable (Tier-1, Tier-2, Tier-3). Institution type was inferred only indirectly from college-name text as an engineering/technology proxy versus other-institution proxy, a limitation highlighted throughout this piece. Screening clearance was a binary outcome only (cleared vs. not cleared).

Primary reason for rejection has been recorded only as one nominal category per person as no multi-item severity index has been created.

7.6 Statistical analysis

We were mainly dealing with ordinal and categorical measurement scale, so we decided to go with non-parametric statistical methods: we have used Mann-Whitney U tests for comparisons of different groups on an ordinal variable and Kruskal-Wallis tests for group comparisons with multiple groups. We have chosen Spearman rank correlation tests to check the relation of ordinal and continuous variables, whereas chi-square test together with Cramér's V and Cliff's delta were our tools to look up categorical relationships and effects of the size. Below we have a Table "Statistical Analysis Plan, " where each hypothesis and its corresponding variables, test(s), and effect-size measure are specified.

Hyp.	Variables	Test(s)	Effect size	Outcome tested
H1	AI skills, proficiency → difficulty	Mann–Whitney U; Kruskal–Wallis; Spearman	rho	Difficulty
H2	Training → confidence	Mann–Whitney U; Spearman	rho	Confidence
H3	City tier → difficulty	Kruskal–Wallis	ϵ^2 (not reported)	Difficulty
H4	AI proficiency, skills → clearance	Chi-square; Cramér's V	Cramér's V	Screening clearance
H5	Primary barrier → confidence	Kruskal–Wallis	ϵ^2 (not reported)	Confidence
H6	Exposure, application count → confidence, difficulty	Mann–Whitney U; Spearman	rho	Confidence; difficulty

H7	Institution-type proxy → confidence, difficulty	Mann–Whitney U / Kruskal–Wallis	not reported	Confidence; difficulty
H8	City tier → technology barrier; technology barrier → confidence	Chi-square; Cramér's V; Cliff's delta	Cramér's V; Cliff's delta	Confidence

Statistical analysis plan by hypothesis.

Due to use of eight hypotheses in multiple bivariate statistics, the analysis may be influenced by a Type I error due to multiple comparisons or the fact of running multiple significance tests without Bonferroni-style correction this was taken into account in the analysis of results as mentioned in Sections 8.5 and Table 6 following Bland & Altman (1995) recommendations. Also, interpretation of the effect sizes throughout relied on the standards of small/medium/large effect benchmarks for the behavioural sciences set by Cohen (1988). According to these benchmarks, the Spearman correlation coefficient, Cramér's V coefficient, and Delta coefficient values obtained in this study consistently fall at or below the limit for a "small effect". None of the source analysis papers performed any multivariate or regression analysis (e. g., ordinal logistic regression of difficulty/confidence or binary logistic regression of screening clearance) was done; in cases where the analysis is enhanced via modelling (e.g., the example of difficulty/confidence being addressed by logistic modelling) this is pointed out in section 13 as a recommendation.

8. Results

Results are presented hypothesis by hypothesis, evidence-first: the test used, the statistic obtained, the decision, and a conservative interpretation. Table 3 reports descriptive comparisons; Table 4 reports the full hypothesis-decision summary; Table 5 isolates effect sizes; Table 6 reports the multiple-testing assessment; Figure 2 reproduces the original descriptive visual comparison for H1–H4.

8.1 Descriptive comparisons

Grouping variable	Difficulty (mean)	Confidence (mean)
AI-related skills — No	2.95	—
AI-related skills — Yes	2.85	—
Formal training — No	—	3.21
Formal training — Yes	—	3.18
City tier — Tier-1	2.85	—
City tier — Tier-2	2.84	—
City tier — Tier-3	3.00	—
AI recruitment exposure — Exposed	2.84	3.10
AI recruitment exposure — Not exposed	2.96	3.29
Institution proxy — Engineering/technology (median)	—	3.0 (median)
Institution proxy — Other (median)	—	4.0 (median)
Barrier categories (all, approx. median)	—	~3.0 (median)

Table 3. Descriptive comparison of perceived difficulty and confidence by grouping variable.

AI proficiency group	Screening clearance rate
Novice	48.2%
Intermediate	50.7%
Advanced	51.8%

Table 3b. Screening clearance rate by AI-proficiency group (H4).

8.2 Hypothesis-by-hypothesis results

H1 — AI skills and perceived difficulty

Participants who indicated skills related to artificial intelligence had a slightly lower mean difficulty score (2.85) than participants who didn't have such skills (2.95), which in principle supported our first hypothesis (H1). However, a Mann-Whitney U test $p=0.531$, Kruskal-Wallis, a level of proficiency, $p=0.444$, and a Spearman correlation $p=0.217$, $\rho=0.063$ did not reveal any statistically significant differences. Therefore, the hypothesis was dismissed. There is no indication in the data that having an artificial intelligence skill or rating your proficiency as higher will lead to a perception of an easier job; the relationship is very weak.

H2 — Formal training and confidence

Mean levels of confidence were hardly different in the trained (3.18) and untrained (3.21) respondents (Mann, Whitney U $p = 0.808$; Spearman $\rho = 0.012$). Decision: not supported.

There was no link of the formal training or certification alone to higher confidence in this group; a valid explanation for that might be that the training quality, the practical exposure, or whether the certification is relevant could be more important than the fact of completion of training this is just one possible interpretation, not a solid finding.

H3 — City tier and perceived difficulty

Mean difficulty scores among Tier-1 (2-85), Tier-2 (2-84), and Tier-3 (3-00) respondents were virtually the same (Kruskal-Wallis $p = 0.601$). Decision: not supported. Although Tier-3 respondents descriptively stated a slightly higher degree of difficulty, the magnitude of the difference and the resulting statistic are neither significant nor sufficient to justify the claim put forth in the hypothetical scenario.

H4 — AI proficiency and screening clearance

There were only small differences in clearance rates depending on levels of proficiency (novice 48.2% intermediate 50.7%, advanced 51.8%), however neither proficiency ($p=0.850$, Cramér's $V = 0.029$) nor AI skills ($p=0.290$, Cramér's $V=0.054$) were significantly different and thus no conclusion can be made. Effect sizes are minute; the AI skills that were self-reported and the levels of proficiency do not seem to adequately predict clearance from the screening in this study.

H5 — Perceived barriers and confidence

The different groups of primary barriers all had the same median belief score of 3.0. They didn't affect each other much (Kruskal-Wallis $p = 0.987$).

Decision: not supported.

Choosing the type of primary barrier didn't depend on confidence, therefore, confidence could be the main determinant rather than a single barrier category. Since only a respondent's main barrier was considered, multiple-choice barrier-severity index would better test this hypothesis in further research.

H6 — Recruitment exposure

Those exposed respondents on the average had a little bit lower confidence (3.10) and a little bit higher difficulty (2.84) when compared to non-exposed respondents with mean confidence of 3.29 and difficulty of 2.96 only.

There was no significant relationship between the exposure and confidence levels as well as the exposure and difficulty (exposure vs. $p = 0.191$ confidence; exposure vs. difficulty $p = 0.401$; application count vs. confidence $\rho = -0.012$, $p = 0.811$; application count vs. difficulty $\rho = 0.051$, $p = 0.315$). Decision: not supported.

The fact that more exposure did not result in greater confidence or lesser difficulty is clear. The only descriptive pattern for confidence was weakly negative but statistically it showed no significance.

H7 — Institution type and readiness

The type-of-institution proxy had a statistically significant relation to the level of confidence ($p = 0.037$): for the engineering/IT proxy group, the median confidence was 3.0, whereas the other-institution proxy group had a median confidence of 4.0. The relationship to difficulty was not statistically significant ($p = 0.578$). Decision: The significance level is only nominal for confidence, but in this case, the analysis results are considered as exploratory since there are two reasons. On the one hand, institution type is not directly measured here, but rather, it is determined based on college names so that there remains the variable-proxy limitations. On the other hand, multiple hypothesis tests are performed at different statistics levels for the same eight hypotheses, i.e. results can be judged against a stringent Bonferroni-corrected threshold of approximately 0.006 (Table 8) and, given a p value of 0.037, the conclusion that it is a false positive has greater statistical plausibility than the alternative.

References: Table 6 & 8 Statistics: t-test, multiple-hypothesis corrections
P-values: significance ($p=0.006$), confidence ($p=0.037$), difficulty ($p=0.578$).

H8 — Technology access barriers and city tier

Technology-access barriers were equally shared across the city tiers ($p = 0.560$), and the participants identifying this obstacle were not significantly less self-confident than the other ones ($p = 0.870$; Cramér's $V = 0.054$; Cliff's $\delta = 0.012$). Conclusion: the hypothesis is not supported.

8.3 Hypothesis decision summary

H	Hypothesized relationship	Test statistic	p-value	Effect size	Decision
H1	AI skills/proficiency → lower difficulty	MWU; KW; Spearman	0.531 / 0.444 / 0.217	$\rho = 0.063$	Not supported
H2	Formal training → higher confidence	MWU; Spearman	0.808 / —	$\rho = -0.012$	Not supported
H3	Tier-1 → lower difficulty than Tier-2/3	Kruskal–Wallis	0.601	not reported	Not supported
H4	Higher proficiency → higher clearance	Chi-square	0.850 / 0.290	$V = 0.029 / 0.054$	Not supported
H5	Barriers → lower confidence	Kruskal–Wallis	0.987	not reported	Not supported
H6	Exposure → higher confidence, lower difficulty	MWU; Spearman	0.191 / 0.401	$\rho = -0.012 / 0.051$	Not supported
H7	Institution type → readiness/difficulty	MWU/KW (proxy)	0.037 / 0.578	not reported	Nominally significant (confidence only); fragile — see Table 6

H8	Technology barrier concentrated in Tier-2/3 → lower confidence	Chi-square; Cliff's delta	0.560 / 0.870	$V = 0.054$; $\delta = -0.012$	Not supported
----	--	---------------------------	---------------	---------------------------------	---------------

Table 4. Hypothesis-testing results and decisions (H1–H8).

8.4 Effect sizes

Table 5 gathers altogether the sizes of the effects presented under various hypotheses. By common measuring rules of the behavioral sciences (Cohen, 1988), almost all numbers indicate effects at most so small even they could almost be negligible. This shows again that even when p-values got closer or crossed the traditional significance lines (as in H7 case), the real impact of the relationship beneath that was still very minor.

Association	Statistic	Value	Interpretation
AI skills/proficiency × difficulty	Spearman rho	0.063	Negligible
Formal training × confidence	Spearman rho	−0.012	Negligible
AI proficiency × screening clearance	Cramér's V	0.029	Negligible
AI skills × screening clearance	Cramér's V	0.054	Negligible-to-weak
Technology barrier × city tier / confidence	Cramér's V	0.054	Negligible-to-weak
Technology barrier × confidence	Cliff's delta	−0.012	Negligible
Application count × confidence	Spearman rho	−0.012	Negligible
Application count × difficulty	Spearman rho	0.051	Negligible

Table 5. Effect sizes and interpretation.

8.5 Multiple-testing assessment

If we consider all the hypotheses (8 total) and more than a dozen individual statistical tests that were conducted, the probability of obtaining at least one false-positive due to chance is actually quite high. Therefore, if we impose a more stringent Bonferroni-corrected cutoff for significance of about 0.006 as done in the original analysis which coincided with the standard Bland and Altman 1995 procedure of a Bonferroni adjustment for multiple testing in one-sided p-value cases), the fact that the one statistically significantly deviating result (H7, $p=0.037$) would still be significant under this adjusted criterion does not hold.

Issue	Nominal p (H7, confidence)	Bonferroni-adjusted threshold	Conclusion
Institution-type proxy vs. confidence	0.037	≈ 0.006 (reported in source material)	Does not survive adjustment — treated as exploratory, not confirmatory

Table 6. Multiple-testing assessment for the nominally significant result.

8.6 Visual comparison

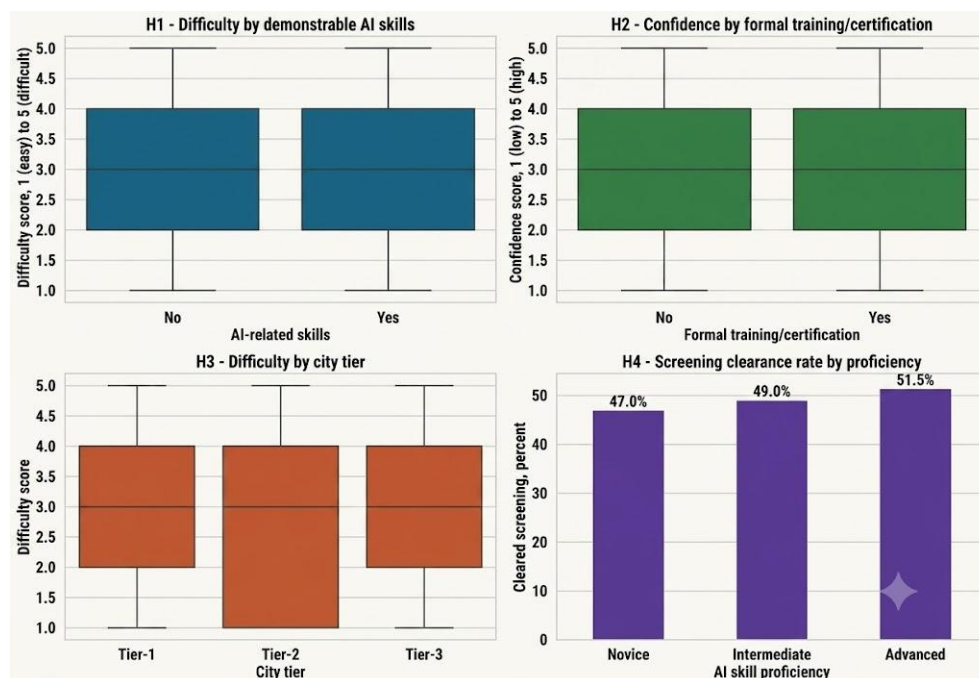


Figure 2. Descriptive comparison of perceived difficulty and confidence by AI-related skills (H1), formal training (H2), city tier (H3), and screening clearance rate by AI-skill proficiency (H4).

8.7 Overall results summary

Among eight hypotheses, six (H1, H6, H8) did not demonstrate a statistically significant relationship at all while the seventh (H7) showed very little change between group medians and the effect sizes were very close to zero in most instances. It was only one hypothesis that yielded a nominally significant result hypothesis 7 (H7), which is however, much lowered by its dependence on a surrogated measure of the variable of interest and a lack of its significance when considering multiple comparisons. In general, the results empirically indicate that possession of AI skills, undergoing formal training programme, working on a city tier, exposure in hiring, types of primary barrier and technology-access barriers were not independent variables for explaining the variance of confidence levels, perceived difficulty levels, and screen success among the people interviewed for this study.

9. Discussion

This study was designed to explore the extent to which the factors widely believed, i. e. AI-related skills, formal training, city tier, institutional context, recruitment exposure, and perceived barriers that lead to preparation are related to the graduates' confidence, difficulty, and selection in recruitment mediated by AI. The main result is predominantly a non-effect pattern: six or eight hypotheses were unsupported whereas the one hypothesis, H7, achieving nominal significance turned out to be unstable in both data collection and testing multiple hypotheses.

The observation that there is mostly no correlation between these variables is again an empirical finding and therefore should not be taken for an unsuccessful study. Theoretical framework of employability (Yorke, 2006) advises explicitly against a simple input-output mechanistic model of human performance enhancement; the outcome of this study clearly agrees with this perspective. It seems very possible that AI-recruitment applicants' readiness is influenced rather by different factors not included in this dataset, e. g. general dispositional self-efficacy, fluency in English in application materials, interview skill in appearance, or even details in the specific design of the AI tools the applicants encountered instead of the isolated skill, training, or situational characteristics that were the focus of testing in this paper.

One way the H1 and H4 negative results can be understood is that self-reported AI skills may be very different in nature to AI competencies that AI screening tools really require. For example, these screening tools might emphasize aspects such as communication, the degree of structure of the answers, or whether the keywords from the resume match the ones used in the AI system, rather than technical AI literacy.

Making sense of the H2 neutral result may be possible if we assume that training completion alone without considering training quality, practical application of training, or perceived legitimacy of certification will not necessarily bring about increased confidence. We must, however, make a very limited claim for the present data since this study does not specify which training aspect(s) can affect the results.

The null H3 and H8 results regarding city tier and technology-access barriers are worth highlighting vis-à-vis digital-divide theory (van Dijk, 2006), that would reasonably guess some level of contextual inequality, in an environment of major urban-rural and tier-based infrastructure disparities. The one reason might be that, these students, who are graduates, might actually already have a baseline level of access to technology beyond city tier such that further variations in access will not influence their views on difficulty or confidence.

This may mean that if a digital divide exists in this sample of people, then it works in some aspects beyond what was captured by city tiers and single-barrier measures in this paper (e.g., motivational/skills-based access, which are different categories of access to technology according to van Dijk), but this paper cannot directly prove that.

H6 result says that exposure to AI-driven recruitment did not lead to higher confidence; in fact, it showed only a very weak, non-significant negative tendency. This finding contradicts an assumed familiarity-breeds-confidence model. One possible explanation is that the use of AI in hiring processes, especially in cases where a candidate got rejected or had to go through a complicated screening procedure, might not only fail to build but also reduce that person's confidence; however, since the present cross-sectional data can't be used to determine one explanatory over the others, including potential reverse associations (e.g., less confident graduates may be applying more or encountering AI-mediated processes more frequently), we can only assume that the exposure to AI in hiring processes could be the cause of the reduction in confidence but we are not sure.

H7 needs to be interpreted extra carefully. Numerically, the difference between the engineering/technology proxy group (median confidence 3.0) and the other-institution proxy group (median confidence 4.0) is the biggest and also gets the most p-value support from all results found in the study. It would be easy to consider that the first thing a paper's about. However, one must refrain from the temptation. As institution type here was determined from college-name texts rather than direct measurement, the grouping variable itself might be misallocating the respondents who fall into these categories. Moreover, since this finding fails to survive a conservative multiplicity correction, it should be seen no more than a promising, but not yet verified, idea of an institutional effect to be checked by other means at a later stage rather than as a solidly established institutional effect.

On the whole, what these findings point to is that AI-recruitment readiness in this sample cannot really be accounted for by the limited number of factors that are, for the most part, easily measurable and therefore have become popular in policy and curricular debates, skill mastery, graduation, and the like. Nevertheless, the absence of an independent linear correlation between such factors and AI-recruitment confidence, difficulty, and screening clearance in this cross-sectional, self-reported, applicant-side dataset implies very little beyond that the factors could still be of great use to graduate employability more broadly.

10. Theoretical Implications

The study results were only loosely in line with the simple interpretation of Capability-to-Outcome model implied by Technology Readiness theory (Parasuraman, 2000), that is, knowing the AI-related skills or just attending the training sessions could not have led a respondent to perceive lower level of task difficulty or, conversely, to have a greater level of self-confidence.

One should not consider, though, Technology Readiness theory is rejected here since the theory primarily describes the personality-related (optimism, inventiveness, anxiety, lack of confidence) components and not the possession-based approach (discrete skills). Therefore in this context it even makes more sense to think about a

personality-based explanation than about a skill-list type explanation. However the dataset did not provide a direct measurement of the TRI-style personality dimensions at all.

Again, the absence of differences related to city-tier and to the presence of technology barriers (H3, H8) does not support the simple physical-access view of the digital divide (van Dijk, 2006), among graduates who took the survey at least, while on the contrary the results conform to van Dijk's own note of caution that digital inequality is multi-faceted and cannot be adequately described with a single tier variable or self-identification of a barrier to a primary access issue.

Considering Yorke's (2006) employability theory, it would be safe to suggest that it is the one explaining, in most cases, the results of the present study: the idea of the workability as a skill, a personality, a context-embedded achievement which is multi-faceted, and not dependent solely on any one factor would perfectly align with what we have observed in the way of dispersed results and many pairs of variables being statistically insignificant.

Hence, the conclusions support AI-recruitment readiness to some degree as multi-factor explanation as opposed to the single-factor ones but remind us as well that to get rid of any doubts or limitations of framework testing one would have to do a multivariate study and use much more detailed, valid measures instead of those that were the available ones of this survey.

11. Practical Implications

11.1 For graduates

A lack of significant relationship between the self-reported AI skills or training of the participants and their confidence or difficulty indicates that graduates need not expect that the completion of a standalone AI-related course or certification will, at least, substantially ease AI-mediated recruitment.

Hands-on experiential learning, to illustrate, participating in a structured simulation of a mock AI screening, could well have a better impact than a piece of paper alone though since the experimental setup here was not quite able to directly prove this the idea is left as a reasonable explanation consistent with the findings but not as a proven fact.

11.2 For higher education institutions

Given the null outcomes of the training and skills measurement, institutions might consider checking if current AI-related training really enhances the specific capabilities and confidence needed to work with AI, rather than assuming that providing any kind of training would be beneficial. Integrating AI-recruitment literacy learning how the tools work, what they are designed to evaluate and how to represent oneself in them as a practical component of a job readiness program may be a better way of addressing job related problems than general AI-skills training.

11.3 For recruiters and AI-recruitment vendors

Weak effects were consistently reported for barrier categories and technology accessibility (H5, H8). Therefore, applicants' confidence might be influenced not only by applicant-side skills or access differences but also by aspects that are under recruiters' control e.g., disclosing the use of AI and the evaluation criteria. Recruiters and vendors can plausibly get motivated to explain AI-screening processes and establish mutual communication channels between them if the job seekers have lower trust in the procedural justice of the AI-based interviewing than that of the human-based one (Acikgoz et al., 2020). Similarly, if the candidates assess a machine-driven selection as less fair compared to human or hybrid decision processes irrespective of whether the outcome works in their favor or not (Lavanchy et al., 2023), then recruiters and vendors have a plausible motive to communicate about their AI-screening processes and explainability. Accepting the evidence that whether or not people agree with automated interview is dependent on the level of involvement (Langer, König, & Papathanasiou, 2019) also suggests that one needs to ensure greater transparency in these situations, where the outcome of the screening is particularly crucial as the one studied in this work.

11.4 For policymakers

There was no association found between city tier and both perceived difficulty as well as technology-access barriers. However, this should not be interpreted as evidence that digital-employability inequality does not exist in India; instead, it suggests that such inequality remained undetected due to this particular sample and these particular measures.

The digital-employability policymakers should consider this not as a sign that spatial digital-access programs are not needed but as a reminder to have more detailed, reliable assessments of digital-employability access.

12. Limitations

12.1 Cross-sectional design

Data were collected at a single time point, precluding causal inference. All reported relationships are associations, not effects.

12.2 Self-reported data

Skills, proficiency, training, exposure, and barrier measures were self-reported and may be subject to recall bias, social-desirability bias, or measurement error in self-assessment; these are discussed as standard methodological possibilities rather than demonstrated biases in this specific sample.

12.3 Institution-type proxy

Institution type was inferred from college-name text rather than measured directly. This is a substantial measurement limitation affecting the interpretation of H7 and should be resolved through direct institutional-type coding in future data collection.

12.4 Primary-barrier measurement

Only the respondent's single primary barrier was analyzed for H5 and H8. A multi-item or multi-select barrier-severity index would likely provide a more sensitive test of barrier-related hypotheses.

12.5 Multiple statistical tests

Eight hypotheses were tested using numerous bivariate statistics, elevating the risk of a Type-I (false-positive) error; consistent with standard guidance on multiple significance testing (Bland & Altman, 1995), this is directly relevant to the interpretation of H7, as detailed in Section 8.5 and Table 6.

12.6 Lack of multivariate modelling

The available analysis is predominantly bivariate and non-parametric. A stronger design would test multiple predictors simultaneously via ordinal logistic regression (for difficulty and confidence) and binary logistic regression (for screening clearance), with appropriately coded control variables such as gender, city tier, institution type, training, and AI skills.

12.7 Generalizability

In the absence of documented sampling procedure, exact sample size, and demographic composition, broad claims about “all Indian graduates” cannot be supported. Findings should be read as applying to the specific, undocumented sample analyzed.

13. Future Research

- Larger and more representative, probability-based samples with documented sampling frames.
- Direct measurement of institution type rather than name-based proxy inference.
- Multi-item barrier-severity measurement in place of single-primary-barrier selection.
- Longitudinal designs tracking graduates across multiple recruitment cycles.

- Objective (employer-verified) screening outcomes rather than self-reported clearance.
- Multivariate modelling, including ordinal logistic regression for confidence and difficulty and binary logistic regression for screening success.
- Interaction effects (for example, training $\times$ city tier) where theoretically justified.
- Cross-city and cross-state comparisons with documented, stratified sampling.
- Direct comparison between AI-mediated and conventional recruitment experiences for the same respondents.
- Validation of a dedicated AI Recruitment Readiness Index if suitable multi-item measurement is developed.
- Direct measurement of Technology Readiness Index dimensions (optimism, innovativeness, discomfort, insecurity) to test dispositional accounts more directly.
- Qualitative follow-up (interviews or open-ended items) to explain why commonly assumed factors showed limited explanatory power.

14. Conclusion

This paper investigated the relationship between AI-related skills, formal training, city tier, institutional context, recruitment exposure, perceived barriers and their impact on confidence, perceived difficulty and screening among Indian graduates who are using AI-mediated recruitment. As the data is mostly ordinal and categorical, non-parametric bivariate tests were used. The analysis showed that none of the seven out of eight researched factors were statistically supported. Moreover, the overall effect sizes were negligible-to-weak. Interestingly, the only statistically significant finding a link between institution-type proxy and confidence after correcting for multiple testing at the conservative level, is a null result based on a proxy method with limitations.

So, the main conclusion from this study is that the individual aspects of preparation commonly accepted to be related to the confidence in AI-mediated recruitment, difficulty in understanding AI, or the success level of the candidates at screenings did not stand out as independent explanations in this sample of data. The report provides the null-dominant pattern of results as a valid empirical contribution without reframing nor minimizing them. It is recommended that research work for the future be guided by larger, more reliable data collections; that the measurement of institutional context be made not by proxies but directly; that barrier and readiness items be presented as multi-item scales; and that multivariate models be employed which will enable the simultaneous testing of several predictors and their interactions. Until solid evidence comes in, the claims of specific, isolated interventions a one-time training course, a listed skill, or a shift in city tier which would consistently make AI-mediated recruitment for Indian graduates easier, should be approached with caution.

References

- [1] Acikgoz, Y., Davison, K. H., Compagnone, M., & Laske, M. (2020). Justice perceptions of artificial intelligence in selection. *International Journal of Selection and Assessment*, 28(4), 399–416. <https://doi.org/10.1111/ijsa.12306>
- [2] Bland, J. M., & Altman, D. G. (1995). Multiple significance tests: The Bonferroni method. *BMJ*, 310(6973), 170. <https://doi.org/10.1136/bmj.310.6973.170>
- [3] Cohen, J. (1988). *Statistical power analysis for the behavioral sciences* (2nd ed.). Lawrence Erlbaum Associates.
- [4] Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly*, 13(3), 319–340. <https://doi.org/10.2307/249008>
- [5] Langer, M., König, C. J., & Papathanasiou, M. (2019). Highly automated job interviews: Acceptance under the influence of stakes. *International Journal of Selection and Assessment*, 27(3), 217–234. <https://doi.org/10.1111/ijsa.12246>

- [6] Lavanchy, M., Reichert, P., Narayanan, J., & Savani, K. (2023). Applicants' fairness perceptions of algorithm-driven hiring procedures. *Journal of Business Ethics*, 188(1), 125–150. <https://doi.org/10.1007/s10551-022-05320-w>
- [7] Parasuraman, A. (2000). Technology readiness index (TRI): A multiple-item scale to measure readiness to embrace new technologies. *Journal of Service Research*, 2(4), 307–320. <https://doi.org/10.1177/109467050024001>
- [8] Raghavan, M., Barocas, S., Kleinberg, J., & Levy, K. (2020). Mitigating bias in algorithmic hiring: Evaluating claims and practices. In *Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency (FAT* '20)* (pp. 469–481). Association for Computing Machinery. <https://doi.org/10.1145/3351095.3372828>
- [9] Suen, H.-Y., Chen, M. Y.-C., & Lu, S.-H. (2019). Does the use of synchrony and artificial intelligence in video interviews affect interview ratings and applicant attitudes? *Computers in Human Behavior*, 98, 93–101. <https://doi.org/10.1016/j.chb.2019.04.012>
- [10] van Dijk, J. A. G. M. (2006). Digital divide research, achievements and shortcomings. *Poetics*, 34(4–5), 221–235. <https://doi.org/10.1016/j.poetic.2006.05.004>
- [11] Wheebox. (2025). *India Skills Report 2025: Global talent mobility — India's decade*. Wheebox, in association with the Confederation of Indian Industry (CII), the All India Council for Technical Education (AICTE), and the Association of Indian Universities (AIU).
- [12] Yorke, M. (2006). *Employability in higher education: What it is – what it is not (Learning and Employability Series 1)*. Higher Education Academy.

Part D — Statistical Integrity Report

Issue	Status	Action taken
Hypothesis reporting (H1–H6, H8)	Reviewed	Reported as statistically not supported; no wording implies significance
H7 nominal significance	Reviewed	Flagged as exploratory; Bonferroni comparison reported explicitly in Results, Discussion, and Table 6
p-values	Verified against source	All p-values reproduced exactly as supplied; none altered or omitted
Effect sizes	Verified against source	Spearman rho, Cramér's V, and Cliff's delta reported verbatim and interpreted as negligible-to-weak
Multiple testing	Addressed	Explicit Bonferroni-adjustment discussion; risk of Type-I error stated in Methodology, Results, and Limitations
Proxy variables	Addressed	Institution-type proxy limitation stated in Methodology, Results, Discussion, and Limitations
Causal language	Screened	“Associated with,” “not supported,” and “differences observed” used throughout; no causal verbs applied to cross-sectional associations
Statistical-test appropriateness	Reviewed	Non-parametric/ordinal-appropriate tests (Mann–Whitney U, Kruskal–Wallis, Spearman, chi-square, Cramér's V, Cliff's delta) judged appropriate for the ordinal/categorical measurement levels used; multivariate/regression modelling flagged as a recommended extension, not fabricated
Sample and methodological detail	Not available	Marked as [INFORMATION NOT AVAILABLE IN SOURCE MATERIAL] rather than estimated or invented
References	Externally verified	All cited works checked against publicly indexed bibliographic records (author, year, venue, volume/pages) before inclusion