

GOVERNANCE DOCS

AI System Impact Assessment

Brightwell Recruitment (sample)

AI screening tool that ranks job applicants for 40 client employers, built on a third-party language model, before recruiters review the shortlist.

ASSESSED	25 September 2026
METHOD	ISO/IEC 42005:2025, for ISO/IEC 42001 clause 6.1.4 and controls A.5.2 to A.5.5
IN SCOPE	8 scope items, 12 risks
APPETITE	Medium and below (level 9 or less) accepted without treatment
PROFILE	Professional services
PREPARED BY	Governance Docs

Confidential. Prepared for Brightwell Recruitment (sample). Contains details of an AI system, its impacts on people and society and the measures that address them; share on a need-to-know basis.

Based on a self-assessment using information supplied by the organization. Not independently verified by Governance Docs; see "Basis of this report and limitation of liability".

CONTENTS

1	Executive summary
2	Screening, the AI system and safeguards
3	Impacts on people and society
4	Measures and residual impact
5	Findings and what closes them
6	Review and decision
A	Criteria used and basis of this report

Where the organization's risks stand

PROCESS SCORE

91 / 100

Level 4 · Managed

Owned, justified and accepted: ready for management review and audit.

Initial

Developing

Defined

Managed

The score measures how complete and defensible the assessment is, not how risky the organization is.

RISKS	ABOVE APPETITE	CRITICAL	ABOVE AFTER TREATMENT	FINDINGS
12	5	0	0	4

WHAT THIS ASSESSMENT TELLS YOU

- Full impact assessment needed: The system has a sensitive use and may be high-risk under the EU AI Act. Assess its impacts in full before deployment. Deployers that are public bodies or provide public services, and deployers of credit scoring or life and health insurance pricing, also owe a fundamental rights impact assessment (Article 27).
- No High or Critical impact remains once the planned measures are in place, so the system can go ahead on those terms.
- 5 of 11 rated risks sit above your appetite line (Medium and below is acceptable).
- The highest risk is R-01, outcomes are worse for some groups of people, at level 16 (High), owned by Head of Data Science.
- Planned treatment brings the number of risks above the line from 5 to 0.

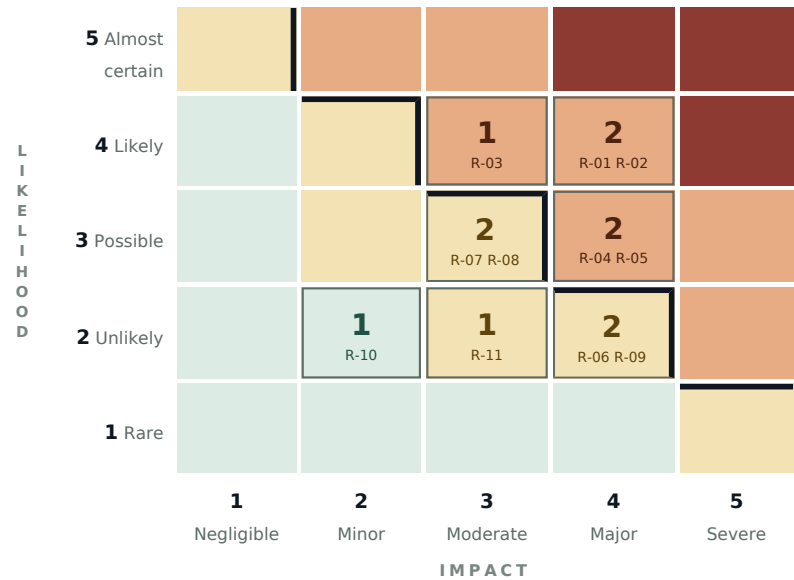
DO THESE FIRST

	SEVERITY	ACTION	RISKS	OWNER	DOCUMENT THAT CLOSES IT
1	MAJOR	Not yet rated for likelihood and impact Rate each one against the likelihood and impact scales in your criteria.	R-12	AI Governance Lead	Risk assessment and treatment methodology
2	MAJOR	Risk with no owner Name an owner, by role, for each risk.	R-11	To assign	Risk assessment and treatment methodology
3	MAJOR	Safeguards missing for some dimensions of impact Put each safeguard in place, or record the impact in the register with the measure that addresses it instead.	People can contest a result and reach a human; Accuracy and fairness are measured before release and monitored in use	To assign	AI system impact assessment

Heat maps: today and after treatment

Each cell shows how many risks sit at that likelihood and impact, with their references. The heavy line is the appetite line: Medium and below (level 9 or less) is accepted without treatment. The second map shows where each risk is expected to be once its treatment is carried out; avoided risks are removed.

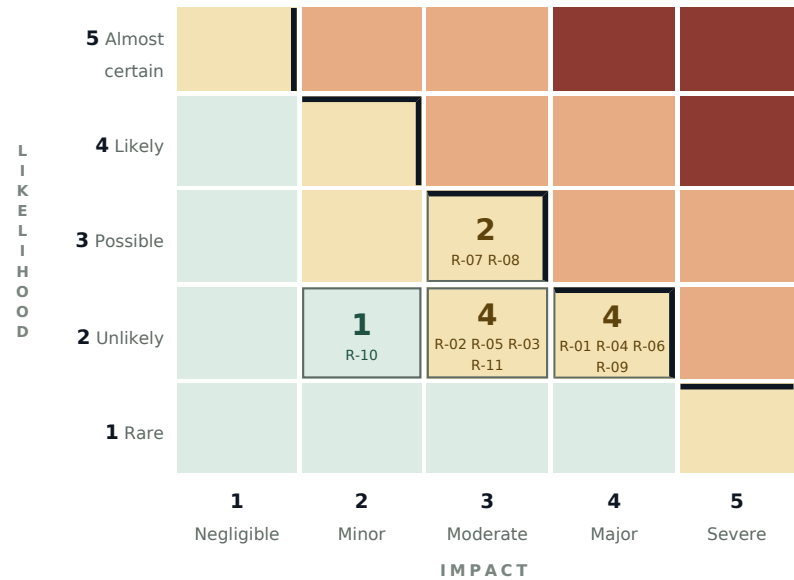
Current risk levels



11 of 12 risks rated.

Low 1-4 Medium 5-9 High 10-16 Critical 17-25 Appetite line

Target levels after treatment



11 risks plotted; 0 removed by avoidance; risks above the line with no decision are not plotted.

RISKS BY LEVEL



Analysis and priorities

IN ONE SENTENCE

A full impact assessment is needed and the planned measures bring every high impact within appetite; add an applicant review route and outcome monitoring before wider launch.

AI-assisted analysis, generated from the information the organization provided and checked automatically against the scores and findings in this report. It supports planning and is not an audit opinion; see "Basis of this report and limitation of liability".

FOR SENIOR MANAGEMENT AND THE BOARD

Our impact assessment of the applicant ranking tool confirms it can go ahead with two clients within the limits we have set, once measures against unfair outcomes are in place. The two gaps to close before any wider launch are a way for applicants to ask for a human review and monthly monitoring of outcomes by group.

WHERE YOU STAND

The screening found a sensitive use, ranking job applicants, that may be high-risk under the EU AI Act, so a full impact assessment is needed before deployment. You have recorded 12 impacts across 8 scope items, and 11 are rated. Your process score is 91 out of 100, at the Managed level.

WHAT IT MEANS FOR THE ORGANIZATION

5 impacts sit above your appetite line, all of them High, and the highest is worse outcomes for older applicants and applicants with career gaps, seen in the pilot. Planned measures bring all five within appetite. Two safeguards are missing: applicants cannot ask for a human review of a ranking, and outcomes are not monitored after launch.

WHERE TO FOCUS FIRST

Put the applicant review route and monthly outcome monitoring in place first, since both are also planned measures for High impacts. Then rate the risk of misleading generated summaries, name an owner for the inference risk and record the owner's acceptance for the drift risk.

PRIORITIES

PRIORITY	WHY IT MATTERS	FIRST STEP	START FROM	WITHIN
1 Give applicants a route to a human review Safeguards missing for some dimensions of impact · R-03 People cannot contest a decision or reach a human	Applicants outside the top 20 are never seen by a person and have no way to challenge the ranking.	Add a review request link to every decision email and agree a 5-day response time.	AI system impact assessment	30 DAYS
2 Monitor outcomes by group after launch R-05 Performance degrades after deployment	New clients and roles differ from the pilot, and the fairness problems it found could return unseen.	Build the monthly dashboard of shortlisting rates by group and record the owner's acceptance.	Risk treatment plan	60 DAYS
3 Rate the risk from generated summaries R-12 Generated content misleads people · Not yet rated for likelihood and impact	Recruiters read an AI-written summary for each applicant, and an unrated impact cannot be compared with your appetite.	Rate likelihood and impact, using pilot examples where summaries were wrong.	Risk assessment and treatment methodology	30 DAYS
4 Name an owner for the inference risk R-11 The system infers sensitive things about people · Risk with no owner	The model could infer age or health from dates and gaps in CVs, and nobody is accountable for checking it.	Assign the risk to the Head of Data Science and add it to the quarterly bias tests.	Risk assessment and treatment methodology	30 DAYS

30 / 60 / 90-DAY ROADMAP

FIRST 30 DAYS • Add the human review route for applicants. • Rate the generated summaries risk and name an owner for the inference risk. • Remove proxy features from the model.	DAYS 31 TO 60 • Launch the monthly outcome dashboard. • Test outcomes by group on the retrained model. • Train recruiters on the tool's limits.	DAYS 61 TO 90 • Gather data to test outcomes for applicants with disabilities. • Review the assessment before adding more clients.
---	--	---

How deep the assessment goes: the screening

Full impact assessment needed

The system has a sensitive use and may be high-risk under the EU AI Act. Assess its impacts in full before deployment. Deployers that are public bodies or provide public services, and deployers of credit scoring or life and health insurance pricing, also owe a fundamental rights impact assessment (Article 27).

Screening note: Recruitment is a sensitive use and is listed as high-risk in Annex III of the EU AI Act. The tool processes applicants' personal data, relies on a third-party model and sits in a process applicants cannot avoid, so a full assessment is needed before launch. Emotion recognition in interviews was considered and ruled out.

Prohibited or restricted uses	Met
The system could fall within a prohibited AI practice	No
Sensitive uses: any one calls for a full assessment	Met
It makes or shapes decisions about people's access to jobs, credit, insurance, housing, education, health care or public services	Yes
It uses biometric data, or identifies or categorises people	No
It could affect people's health, safety or physical wellbeing	No
It affects children or other vulnerable people	No
It is classed as high-risk under the EU AI Act or a similar law	Yes

Context: two or more call for a full assessment	Met
It affects many people	Yes
Its outputs take effect without meaningful human review	No
It interacts directly with people, or generates content they see	No
It uses personal data	Yes
It relies on a model you did not build and cannot fully inspect	Yes
People would find it hard to avoid or opt out of it	Yes

The AI system

Following ISO/IEC 42005: the system and its purpose, intended and unintended uses, data, model, deployment environment, the people affected and the expected benefits. The systems, data, models and groups involved are listed with the register.

The AI system and its purpose	The tool reads CVs and application answers, scores each applicant against the job requirements and ranks them. Recruiters see the top 20 per role with a short summary of why each was ranked.
Intended uses and users	Used by Brightwell recruiters for client roles in the UK, Ireland and Germany; about 60,000 applications a year. Recruiters make every shortlist and hiring decision; the tool only ranks.
Unintended uses and reasonably foreseeable misuse	Clients asking for the ranking to reject applicants automatically; recruiters shortlisting only from the top 20 under time pressure; applicants using AI-written CVs to game the score.
Data: sources, quality and representativeness	Fine-tuned on 5 years of Brightwell placements (about 180,000 applications and outcomes). Past placements under-represent applicants over 50 and applicants with career gaps.
Model and algorithm	A general-purpose language model from an AI provider, fine-tuned by Brightwell and called through the provider's API. The provider gives a model card but not its training data.
Deployment environment and human oversight	Cloud service hosted in the EU. Recruiters can see the reasons for a ranking and can move any applicant onto the shortlist. The system can be switched off per client.
People, groups and society affected	Applicants, including older applicants, applicants with disabilities and those with career gaps; recruiters, whose work changes; client employers.
Expected benefits	Applicants get a reply within 48 hours instead of 2 weeks; recruiters spend time on interviews instead of first sifts; more consistent criteria across clients.
Lifecycle stage and trigger for this assessment	Before deployment, after a 3-month pilot with two clients.

Safeguards for each dimension of impact

Whether the safeguards that limit harm are in place, one question per dimension of impact. Every "partly" or "no" says what is missing.

Question	Answer	Explanation or evidence
Outcomes have been tested for unfair differences across the groups affected	Partly	Pilot tested by sex and age; not yet by disability or ethnicity, where data is sparse.
People are told when AI is involved, and content it generates is labelled	Yes	Application page explains that AI ranks applications and a recruiter reviews them.
People can get a meaningful explanation of results that affect them	Partly	Recruiters see reasons; applicants get a general explanation only, on request.
A competent person can review, override or stop the system's outputs	Yes	Recruiters can add any applicant to the shortlist; weekly review of overrides.
People can contest a result and reach a human	No	No route for applicants to ask for a human review of a ranking.
Personal data use is lawful and minimised, with a DPIA where one is needed	Yes	DPIA completed; legitimate interests assessment recorded; the provider does not train on our data.
Failures and foreseeable misuse have been tested, with a safe fallback	Partly	Tested against AI-written CVs; no fallback if the provider API fails.
Accuracy and fairness are measured before release and monitored in use	No	Accuracy measured in the pilot only; no ongoing monitoring planned.
Wider effects on society, jobs and the environment have been considered	Yes	Effect on recruiters discussed with the team; no job losses planned.
Responsibilities for the system, its monitoring and its incidents are assigned	Yes	AI Governance Lead owns the system; incidents go through the complaints process.

The actual and reasonably foreseeable impacts on individuals, groups and society follow in the register, with the measures that address them and their ISO/IEC 42001 Annex A references.

What was in scope

The AI system or feature, use or decision, people or groups affected, data, model and AI provider or partner the assessment covered, with their owners.

Type	Scope item	Owner	Risks
AI system or feature	Applicant ranking tool	AI Governance Lead	4
Use or decision	Shortlisting for client roles	Head of Recruitment Operations	2
People or groups affected	Job applicants	Head of Recruitment Operations	3
People or groups affected	Applicants over 50 and with career gaps	Head of Recruitment Operations	1
People or groups affected	Recruiters	HR Director	2
Data	Historical placement records	Head of Data Science	1
Model	Fine-tuned language model	Head of Data Science	4
AI provider or partner	AI model provider	Procurement Manager	1

Risk register

Every risk, highest level first. Level = likelihood x impact, rated with the controls already in place. The rationale is the organization's own reason for the rating. The full scenario for each risk (threat, weakness and consequence) is in the workbook.

#	Ref	Risk scenario	Scope items	Dimensions of impact	Owner	Existing controls	L	I	Level	Rationale
1	R-01	Outcomes are worse for some groups of people Weakness: Not tested for unequal outcomes across the groups it affects, before or after release	Applicant ranking tool, Applicants over 50 and with career gaps	Fairness, Rights and access	Head of Data Science	Pilot tested by sex and age	4	4	16 · High	In the pilot, applicants over 50 were shortlisted at half the rate of others with similar experience.
2	R-02	Historical data carries past discrimination into new decisions Weakness: Training data not reviewed for bias or representativeness	Historical placement records, Fine-tuned language model	Fairness, Rights and access	Head of Data Science		4	4	16 · High	Five years of past placements reflect past client preferences, including against career gaps.
3	R-04	Staff rubber-stamp the system's outputs Weakness: Oversight exists on paper, with no training, time or authority to override	Shortlisting for client roles, Recruiters	Fairness, Autonomy and oversight	Head of Recruitment Operations	Recruiters can add any applicant	3	4	12 · High	Recruiters added applicants from outside the top 20 in fewer than 2% of pilot roles.
4	R-05	Performance degrades after deployment Weakness: No monitoring of accuracy or fairness in use	Applicant ranking tool, Fine-tuned language model	Fairness, Safety and wellbeing	Head of Data Science		3	4	12 · High	New clients and roles differ from the pilot; there is no monitoring after launch.
5	R-03	People cannot contest a decision or reach a human Weakness: No appeal route and no human review on request	Shortlisting for client roles, Job applicants	Autonomy and oversight, Rights and access	Head of Recruitment Operations	General complaints process	4	3	12 · High	Applicants ranked outside the top 20 are never seen by a person and cannot ask for a review.

6	R-07	People cannot get an explanation of a result that affects them Weakness: The logic is opaque and staff cannot explain individual outputs	Job applicants, Applicant ranking tool	Transparency, Rights and access	Head of Recruitment Operations	General explanation on request	3	3	9 • Medium	Applicants get a standard paragraph, not the reasons for their own ranking.
7	R-08	The provider's model behaves in ways you cannot see or control Weakness: No information from the provider on testing, limits or changes	AI model provider, Fine-tuned language model	Transparency, Safety and wellbeing	Procurement Manager	Model card reviewed	3	3	9 • Medium	The provider updates the base model without notice; a pilot update changed rankings for 6% of applicants.
8	R-06	Some people cannot use the system or are shut out by it Weakness: Accessibility not tested; no alternative route to a person	Job applicants	Fairness, Rights and access	Head of Recruitment Operations	Accessible application form	2	4	8 • Medium	Applicants who declare a disability are few in the data, so outcomes for them cannot yet be tested.
9	R-09	The system is misused in a reasonably foreseeable way Weakness: Foreseeable misuse not identified; no usage limits or monitoring	Applicant ranking tool	Safety and wellbeing, Society and environment	AI Governance Lead	Client contract forbids automatic rejection	2	4	8 • Medium	Two clients asked during the pilot whether the ranking could reject applicants automatically.
10	R-11	The system infers sensitive things about people Weakness: Nobody checked what the model can infer, or who sees it	Fine-tuned language model	Privacy, Rights and access	None		2	3	6 • Medium	The model could infer age or health from gaps and dates in CVs.
11	R-10	Jobs and working conditions change for the worse Weakness: Effects on staff not assessed and staff not consulted	Recruiters	Autonomy and oversight, Society and environment	HR Director	Team briefing	2	2	4 • Low	No job losses planned; first-sift work moves to interviews.
12	R-12	Generated content misleads people Weakness: No labelling, no source links and no check on accuracy	None linked	Transparency, Safety and wellbeing	AI Governance Lead		-	-	Not rated	

What will be done, by whom and by when

Every risk with a treatment decision, earliest due date first. The iso 42001 annex a controls listed are the ones the treatment relies on; they carry into the ai system impact assessment. Acceptance is the risk owner's approval of the plan and of the risk that remains.

Ref	Risk	Decision	Planned actions	ISO 42001 Annex A controls	Owner and due date	Current to target	Accepted
R-01	Outcomes are worse for some groups of people Head of Data Science	Modify	Remove proxy features (graduation year, gap length), rebalance training data, test outcomes by group before launch and every quarter.	A.7.4 Quality of data for AI systems A.6.2.4 AI system verification and validation A.5.4 Assessing AI system impact on individuals or groups	Head of Data Science 30 Nov 2026	16 · High → 8 · Medium	Managing Director 22 Sep 2026
R-02	Historical data carries past discrimination into new decisions Head of Data Science	Modify	Review training data for representativeness and drop outcomes from clients with known skewed hiring.	A.7.4 Quality of data for AI systems A.7.5 Data provenance A.7.6 Data preparation	Head of Data Science 30 Nov 2026	16 · High → 6 · Medium	Managing Director 22 Sep 2026
R-03	People cannot contest a decision or reach a human Head of Recruitment Operations	Modify	Add a "request a human review" link to every decision email, answered within 5 working days.	A.9.2 Processes for responsible use of AI systems A.8.5 Information for interested parties A.3.3 Reporting of concerns	Head of Recruitment Operations 15 Dec 2026	12 · High → 6 · Medium	Managing Director 22 Sep 2026
R-04	Staff rubber-stamp the system's outputs Head of Recruitment Operations	Modify	Show recruiters a sample from below the cut-off on every role, and train them on the tool's limits.	A.9.2 Processes for responsible use of AI systems A.4.6 Human resources	Head of Recruitment Operations 31 Jan 2027	12 · High → 8 · Medium	Managing Director 22 Sep 2026

R-05	Performance degrades after deployment Head of Data Science	Modify	Monthly dashboard of shortlisting rates by group and by client, with an alert threshold.	A.6.2.6 AI system operation and monitoring A.6.2.8 AI system recording of event logs	Head of Data Science 31 Jan 2027	12 · High → 6 · Medium	Not recorded
------	--	---------------	--	---	-------------------------------------	------------------------	--------------

6 risks are within appetite (Medium and below) and need no treatment; 1 are not yet rated.

Where the assessment falls short

Each finding is a gap between the assessment as recorded and what ISO 42005 clauses 5 and 6 ask for, with the risks it applies to and the document that closes it.

Severity	Finding	Applies to	What it means	To fix	Document that closes it
MAJOR	Not yet rated for likelihood and impact	R-12	Without both ratings a risk has no level, so it cannot be compared with the appetite line or placed in priority order.	Rate each one against the likelihood and impact scales in your criteria.	Risk assessment and treatment methodology
MAJOR	Risk with no owner	R-11	Each impact needs an owner who can put the measures in place. Without one, the assessment records an impact nobody will reduce.	Name an owner, by role, for each risk.	Risk assessment and treatment methodology
MAJOR	Safeguards missing for some dimensions of impact	People can contest a result and reach a human; Accuracy and fairness are measured before release and monitored in use	Each "no" leaves a dimension of impact without a safeguard. ISO/IEC 42001 controls A.5.4 and A.5.5 expect impacts on individuals, groups and society to be assessed and addressed.	Put each safeguard in place, or record the impact in the register with the measure that addresses it instead.	AI system impact assessment
MINOR	Residual risk not yet accepted by the owner	R-05	The system owner should approve the measures and accept the impact that remains. If high impact remains, the system should not be deployed as it stands.	Record the owner's acceptance once they have approved the plan.	Risk treatment plan

PROCESS SCORE IN DETAIL

Criteria defined Every likelihood and impact level has a definition, and the appetite line is set.		100%	11 of 11	weight 10
Scope covered Every scope item has at least one risk linked to it.		100%	8 of 8	weight 10
Risk owners named Every risk has an owner accountable for it.		92%	11 of 12	weight 15
Risks rated Every risk has a likelihood and an impact, so it has a level.		92%	11 of 12	weight 15
Ratings explained Existing controls and a rationale are recorded for each risk.		79%	19 of 24	weight 10
Treatment decided Every risk above the appetite line has a treatment option chosen.		100%	5 of 5	weight 15
Targets within appetite Treatments bring each risk to a target at or below the line.		100%	5 of 5	weight 10
Treatments scheduled Each treatment has an owner and a due date.		100%	5 of 5	weight 5
Residual risk accepted Risk owners have accepted what remains after treatment, by name and date.		80%	4 of 5	weight 10
Impact assessment record complete Screening, the AI system description, the safeguards, the governance review, the views of those affected, the decision and the review date are all recorded.		83%	19 of 23	weight 25

Findings by severity



Critical

Major

Minor

- 1 • Initial (0+)**
A start on a register, not yet an assessment that decisions can rest on.
- 2 • Developing (40+)**
The risks are identified, but gaps in rating or treatment would not survive an audit.
- 3 • Defined (60+)**
A complete, repeatable assessment with a treatment plan behind it.
- 4 • Managed (80+)**
Owned, justified and accepted: ready for management review and audit.

The decision on the AI system

RESIDUAL IMPACT: CAN THE SYSTEM GO AHEAD?

No High or Critical impact remains once the planned measures are in place, so the system can go ahead on those terms.

GOVERNANCE REVIEW

Who advised	Reviewed by an AI governance committee or ethics board
Advice	AI governance committee: do not launch until outcomes are tested by disability and an applicant review route exists; monitor ranking outcomes by group every quarter.
Followed	Partly
Why not in full	Disability testing needs more data; launch with two clients while it is gathered, with recruiters reviewing every application from applicants who declare a disability.

VIEWS OF THE PEOPLE AFFECTED

Views	Sought, and summarised below
Summary or reason	Pilot applicants surveyed (412 replies): 71% preferred faster replies; 38% wanted to know how the ranking worked; several older applicants were concerned about fairness.

OUTCOME AND APPROVAL

Outcome	Proceed only within stated limits on its use
Approved by	Managing Director
Approved on	22 Sep 2026
Next review by	31 Mar 2027

Review the assessment when the system changes in a way that could change its impacts (a new use, new data, a new model, new users or a new country) and at the date above. If an unacceptable impact remains that no measure can reduce, the system should not be deployed.

Criteria used and basis of this report

An AI system impact assessment looks at the consequences of the system for individuals, groups and society, so likelihood and severity are rated for them, not for the organization. This is the record of the criteria every rating in this report was made against.

LIKELIHOOD SCALE

Value	Label	Definition
1	Rare	Not expected in the next 5 years; no known case here or in the sector
2	Unlikely	Could happen once in 2 to 5 years; has happened in the sector
3	Possible	Could happen about once a year
4	Likely	Expected several times a year
5	Almost certain	Expected monthly or more often, or already happening

IMPACT SCALE

Value	Label	Definition
1	Negligible	Negligible: people are not affected, or only trivially
2	Minor	Minor: inconvenience people overcome without difficulty
3	Moderate	Significant: unfair or wrong outcomes people can put right with some difficulty
4	Major	Serious: discrimination, lost opportunities, financial loss or distress
5	Severe	Severe: serious, lasting or irreversible harm to health, safety, livelihood or rights

RISK LEVELS AND ACCEPTANCE

Band	Levels	Treatment
Low	1-4	Accepted without treatment
Medium	5-9	Accepted without treatment
High	10-16	Needs a treatment decision
Critical	17-25	Needs a treatment decision

Level = likelihood x impact. A risk above the appetite line may still be retained, but only with the risk owner's recorded acceptance.

Process score. Ten dimensions: nine for the register, each the share of impacts (or of impacts above the line) that meet it, and one for the completeness of the impact assessment record. Each is multiplied by its weight and the total is scaled to 100. A dimension with nothing in scope counts in full. Any critical finding holds the level at Developing, whatever the score.

Your workbook. The Excel workbook that comes with this report is the register as a working file. Change a likelihood or impact on the Risk register sheet and every level, the heat maps on the Dashboard and the process score recalculate. It also holds the treatment plan and a list of the iso 42001 annex a controls your treatments rely on.

AI impact assessment for Brightwell Recruitment (sample), generated 25 September 2026 by Governance Docs. This is a structured self-assessment based on the organization's own criteria and ratings. It is not a certification, an audit opinion or a substitute for an independent risk assessment.

BASIS OF THIS REPORT AND LIMITATION OF LIABILITY

This report is the output of a self-assessment. The risks, ratings, risk levels, treatment decisions, findings, analysis and recommendations in this report are derived solely from the information entered by the organization, and have not been independently reviewed, verified, tested or audited by Governance Docs. Governance Docs did not conduct an independent assessment of the organization, its controls or its operations.

The accuracy and completeness of the results depend entirely on the accuracy and completeness of the information provided. Governance Docs gives no assurance, certification or audit opinion, and accepts no liability for any inaccuracy, omission or deviation arising from that information, or for any decision taken, or not taken, in reliance on this report.

The report, including any AI-assisted analysis, is provided for guidance and planning only. It does not constitute legal, regulatory or other professional advice, and it does not replace an independent assessment or a certification audit.