

AI TASK EXPOSURE AND WORKFORCE PLANNING

Vega AiShift

Structures a discussion of how AI may affect tasks in a role. The next organisation-level training-needs workflow is being developed separately.

METHODOLOGY AND DEVELOPMENT DIRECTION

Which tasks warrant AI experimentation, human oversight or further investigation?

Analyse tasks within a role

The documented role-level method takes a job description, CV or syllabus and identifies tasks. Its unit of analysis is the task. The output is not an assessment of an individual's performance or a forecast of their employment.

Estimate exposure and ground the match

A fixed rubric guides model-based exposure judgements. Tasks are matched to the O*NET occupational catalogue, with available catalogue ratings contributing context. The occupational match and the model judgement are different parts of the evidence.

Allow an undetermined result

A weak catalogue match can trigger abstention instead of a score. An undetermined task means the method lacks sufficient grounding, not that the task is protected from AI or unimportant to the business.

Move towards organisational context

The September programme calls for cumulative organisation-level training-needs analysis across business context and source revisions. The available interaction prototype is synthetic and unscored; it is not evidence that the new diagnostic framework is approved or live.

METHOD FIRST. INTERPRETATION AND EVIDENCE LIMITS ON PAGE 2.

READING THE RESULT

What the evidence supports

Review the underlying task description

A verdict is only as useful as the task that was extracted and matched. Domain experts should check whether the task is accurate, whether its context is represented and what human accountability remains.

Treat benefits as hypotheses

The source paper identifies unsourced assumptions behind recoverable-hours estimates and unmeasured model variation. Neither estimated time savings nor a category label should be presented as a guaranteed financial return.

Keep planning separate from decisions about people

AI exposure is not employee performance. Use the output to design experiments and learning discussions, with people responsible for checking feasibility, privacy, safety and the actual quality of the resulting work.

Evidence limits

The source paper states that expert-agreement and outcome validation had not been completed and run-to-run model variance was unmeasured. Its pending unrated-task weighting correction is now implemented in the current code. The organisation-level programme still requires an approved analytical framework and agreed pilot terms before diagnostic claims can be made.

Appropriate use

Use the documented method as a structured starting point for workforce planning or learning and development review. Do not use it alone for redundancy, employee evaluation or claims about realised productivity gains.

A practical next step

Select a task for a supervised, measurable experiment and agree success criteria before using AI. For the organisation-level workflow, first agree the evidence mappings, uncertainty rules, source permissions and review process.

Source: Vega AiShift methodology white paper, 11 September 2026. This concise adaptation preserves the source's evidence limitations. It is not a fresh validation study or a change to an assessment release. Product availability and programme direction should be confirmed before a deployment.