

Claude occupation evidence: data and methods

Source definitions, statistical support and reproducible measurement

Fatih Kansoy | 9 September 2026

This data note explains how public Claude activity becomes an occupational dataset and why its layers must be distinguished. The purpose is to support reproducible analysis of published occupation shares, task breadth and classification sensitivity for Europe and supported comparators, while retaining all possible source-supported geographies in a global inventory.

The complete registry has 250 entries. Direct occupation evidence is available for 121 May geographies; the task archive spans 581 country-periods. The note documents primary-source fields, formulas, statistical joins and the empirical findings that determine the appropriate use of those files.

The contribution is an observable set of measurement choices and checks. These include country volume and per-capita metrics, a direct occupation panel, typed mapping alternatives, native US employment, theoretical comparisons and partial survey content validation. The paper develops the motivation and literature; this note supplies the technical route and quantitative evidence.

The observations are occupation-tagged consumer Claude conversations. They do not identify the user's actual job, workforce adoption, hours automated, productivity gains or causal employment effects.

1. Research question, motivation and contribution

Generative AI is used for activities associated with work, but a conversation about a task is not a worker, a job or an hour of production. This distinction becomes consequential when platform records are joined to occupational statistics. A country may publish more task categories simply because it contributes more conversations, while an occupational crosswalk can alter the apparent distribution even when the underlying observations are held fixed.

This study asks which occupational descriptions of published Claude use are supported by the public data, how publication and classification choices affect those descriptions, and where official employment and independent activity surveys help interpret them. The substantive focus is Europe, including the United Kingdom and Türkiye, with four supported comparators. A global inventory retains every source-supported geography and records the remaining gaps.

The contribution has three parts. First, we distinguish country usage volume, per-capita platform usage, direct occupation shares and task coverage, and quantify their different publication support. Second, we make typed occupational mappings and their alternatives inspectable, while retaining a native US occupation route. Third, we place the resulting measures beside Anthropic exposure, theoretical capability scores and PIAAC work activities, reporting both coefficients and comparison support.

The purpose is measurement discipline before labour-market inference. The dataset provides reproducible descriptions and diagnostics that can inform later research designs. It does not claim to be the first international platform-use dataset or an exact European reconstruction of Anthropic's exposure score.

2. Platform use and geographic evidence

Handa et al. (2025) establish the task-based analysis of Claude conversations underlying the Economic Index. Chatterji et al. (2025) study consumer ChatGPT use, including the distinction between work and non-work activity. Tomlinson et al. (2025) use Copilot interactions to study occupational applicability. Together these studies show the value of directly observed interactions and the need to distinguish tasks represented in use from people's occupational identities.

The closest geographic antecedent is Anthropic's own work. Appel et al. (2025) describe uneven country and enterprise use and introduce a usage index relative to working-age population. The June 2026 Cadences release supplies monthly country usage shares, a per-capita index and occupation as well as task classifications. We use those published country fields directly. A task-cell count is not a substitute for the publisher's usage-volume or per-capita series.

Fan and Nguyen (2026) combine usage evidence with labour costs to study the distribution of potential gains. That valuation exercise motivates attention to occupation and country denominators. Our analysis addresses an earlier measurement question: how much information survives publication and mapping, and how closely do the resulting occupation descriptors align with other constructs? Employment-weighted taxonomy means are not labour-cost-equivalent gains.

Steele and Cruz (2026) compare career-related AI projections and explicitly consider disagreement across models. The present study treats disagreement between mapping routes and measurement families as evidence to report, rather than collapsing different quantities into a single apparently precise score.

3. Exposure, outcomes and the UK context

Eloundou et al. (2023) assess whether language models can reduce the time required for tasks. Felten, Raj and Seamans (2021, 2023) relate AI capabilities, including language modelling, to occupational abilities. The ILO (2025) and OECD measures similarly concern technological exposure. These are capability constructs: a task may be feasible yet rarely observed on a particular platform. Rank agreement with them is useful construct evidence, but is not a validation against observed worker adoption.

Brynjolfsson, Li and Raymond (2025) study productivity in a concrete workplace deployment. Humlum and Vestergaard (2025) connect adoption evidence to labour-market outcomes in Denmark, while Bick, Blandin and Deming (2024, revised 2025) measure adoption using population surveys. Their designs underscore the additional ingredients needed for claims about workers and outcomes: an identified population, measures of actual use and an outcome design. Public conversation shares alone do not supply these ingredients.

Two UK applications, the Centre for British Progress report (2026) and the Bank Underground article (2026), illustrate policy interest in occupational AI measures. We examine their public documentation and recoverable chart data. Missing composite weights, exact mappings and outcome panels prevent a controlled reconstruction of their full pipelines. We do not infer that either used the particular untyped crosswalk included here.

The connection to this literature is therefore specific. We contribute a documented bridge from published platform classifications to occupational descriptions, an empirical account of publication support, and checks on the interpretation of the bridge. We make no causal claim about hiring, displacement or earnings from the descriptive comparisons that follow.

4. Measurement framework

Let $U(c,m)$ be a country's percentage share of global consumer usage and $A(c,m)$ the published usage-per-capita index. The latter divides usage share by the country's share of population aged 15-64. It is a relative platform-use index, not a fraction of residents who use AI. Both retain the source denominator and two-decimal rounding.

The primary occupation descriptor $S(c,m,o)$ is the direct detailed-occupation share published in the `soc_occupation` facet, divided by 100. The dataset retains the raw source value when a row exists. It separately defines published occupation mass as zero when a row is absent, with actual unreported usage marked unknown. Published rounded-zero rows and absent rows are distinguishable. A zero in the published-mass variable is not an imputed zero in actual use.

The secondary descriptor $C(c,m,o)$ is the fraction of the fixed O*NET task catalogue $T(o)$ with positive published task cells. Its companion $I(c,m,o)$ is the sum of those task shares, in fractions, divided by the catalogue task count. Occupation membership can overlap across tasks. Neither the sum of C nor the sum of I is a national usage total.

When S passes through ESCO, the output is an equal mean of linked donor shares, followed by equal ESCO means within ISCO4. This semantic donor mean is not an allocation of a national conversation total. Source occupation shares are kept available in their native classification so users can distinguish a published share from its mapped descriptor.

Quantity	Denominator	Interpretation
U: global usage share (%)	All global consumer usage in the source window	Relative country volume
A: per-capita usage index	Country share of population aged 15-64	Relative platform intensity
S: direct occupation share	Country consumer usage	Published occupation-tagged mass
C: task coverage	Tasks in the pinned occupation catalogue	Breadth of positive publication
I: task-share intensity	Catalogue task count	Published usage-share fraction per task
Mapped S	Available semantic donors	Taxonomy mean, not probability allocation

Source: Anthropic June data documentation; definitions implemented in `occupation_evidence.py` and `build_country_panel.py`.

5. Data and observation support

The pinned sources produce a registry of 250 countries and areas, 180 with country source rows, and 128 with positive task evidence across 581 country-periods. The fixed task catalogue is O*NET 30.2. All geography labels and source vintages are retained; the registry is not a claim of universal measurement.

Three one-week snapshots in August 2025, November 2025 and February 2026 are retained as task diagnostics. April and May 2026 are calendar-month observations with the direct occupation and overall-usage fields used here. Weekly and monthly windows differ in sampling and classification and are not a harmonized adoption time series. The modern consumer release covers chat and Cowork under its documented plans; first-party API data have no country dimension and do not enter country estimates.

The May task sample contains 121 geographies. 70 have fewer than 100 positive task cells and 85 have fewer than 200. The complete inventory keeps sparse countries available for appropriately limited analysis. Country-level overall fields are not fabricated where absent.

The June appendix (pages 2 and 10) documents hourly sampling, a two-step task classifier, randomized selection among highest-confidence candidate tasks, and examples of classification error. Occupational tags concern the activity assigned by the classifier; even a fully published facet is not a verified record of the user's job.

Employment is selected from official Eurostat, ILOSTAT and UK APS tables at the finest eligible classification detail, with source-specific years, population, age and ICLS definitions. The US uses BLS base-year employment directly. Older or coarse employment observations remain visible and lower the reporting support grade; they are not relabelled as current detailed statistics.

Source schema, keys and missingness

occupation_evidence.py selects modern country rows with category overall and metrics usage_pct / usage_per_capita_index. It separately selects category soc_occupation, hierarchy level 0, metric pct and valid detailed O*NET codes. Hierarchy level 0 in the onet facet denotes task detail; it is not the occupation facet. Unique country, period and occupation keys are asserted before joins.

The raw occupation value, published-row flag, positive flag, rounded-zero flag and source filename remain attached. The full catalogue join adds occupations with no published source row only within supported monthly country windows. Their raw value remains missing, their published mass is zero and their actual usage status is unknown. Overall country rows outside the task sample remain available without generating fabricated task evidence.

Shares retain the country consumer denominator, with no rescaling to released rows. Facet totals count each source occupation or task once, before mapping expansion. A task that belongs to several occupations can contribute to each occupation descriptor but never multiplies the country total. Reported occupation-mass rounding ranges allow 0.005 percentage points per numeric row and truncate to [0,100]; they are conservative arithmetic ranges, not confidence intervals.

The source snapshot is pinned to the documented Hugging Face revision. Release directories preserve publisher vintages, and the O*NET catalogue is fixed at 30.2. A build against a newer unpinned taxonomy or an unreviewed publisher schema is not asserted to reproduce this dataset.

6. Published task breadth tracks usage volume

Across the 121 May geographies, the Spearman correlation between positive published task-cell counts and the country share of global usage is 0.9960. Regressing log task-cell counts on log global usage share yields slope 1.034 and R squared 0.935. The relationship is also strong in April. These are empirical publication-volume diagnostics, not an assertion that the volume mechanism is wholly unmeasurable.

The regression has an intercept and uses positive, nonmissing source shares and task counts. There are no excluded zero or missing usage shares among the task countries in either displayed sample. HC3 standard errors are shown for transparency; their usual sampling interpretation is limited because geographies enter through source availability and publication rules, rather than a random draw of countries.

The near-unit elasticity shows why published task breadth is a poor standalone indicator for comparing country adoption. More heavily represented geographies reveal more categories. However, the regression does not separately identify sample size, disclosure thresholds, classifier coverage and genuine task composition. Unrounded conversation counts and the unpublished distribution are unavailable. It is therefore too strong to equate every additional task cell with volume alone or to infer a causal effect from this association.

For country-level platform intensity, the release supplies U and A explicitly. For occupation analysis, the next result motivates using direct occupation shares as the primary published descriptor and task coverage as a secondary breadth measure.

Window	Countries	Spearman	Log-log slope	HC3 SE	R squared
2026-04-01	114	0.992	1.121	0.051	0.937
2026-05-01	121	0.996	1.034	0.047	0.935

Table 1. Natural-log OLS with an intercept. Source: outputs/evidence/volume_regression.csv. R squared and the slope are descriptive, not causal estimates.

7. Direct occupation shares retain more published mass

The median May sum of published task shares is 42.25% of country usage, compared with 79.17% for the direct occupation facet. Thus the choice of facet materially changes the available occupational evidence. The table reports source-share sums before any occupation crosswalk or employment weighting.

In the United Kingdom, 437 detailed occupations have positive direct occupation shares, while only 324 catalogue occupations have any positive published task. The occupation facet is therefore useful even when within-occupation task detail is thin. This motivates the primary descriptor; it does not eliminate platform selection.

These sums measure retained published mass. The residual combines unpublished or unclassified activity and rounding; it is not all attributable to privacy suppression. Different facets also classify the same underlying activity at different levels. The mass contrast is not a recovery of the missing task distribution.

No share is renormalized to the set of published categories. The machine-readable outputs include numeric-row counts, rounded zeros and conservative rounding ranges. The extra occupation detail cannot reveal users' jobs or representativeness, and a coarse facet can retain more mass while conveying less information about task breadth.

Country	Positive tasks	Task mass (%)	Occupation mass (%)	Positive occupations
United States	1104	92.44	98.29	515
United Kingdom	680	84.96	97.67	437
Netherlands	353	77.18	95.64	307
Türkiye	290	71.69	94.61	281
Median of 121	63	42.25	79.17	109

Table 2. May 2026, original country consumer denominators. Source: outputs/evidence/facet_comparison.csv. Positive occupation counts exclude displayed zeros.

8. Geographic focus and support grades

Support grades combine publication breadth, retained occupation mass, repeated monthly observation, employment recency, statistical detail and matched employment. They are transparent reporting rules adopted after inspection of the public data, not a preregistered sample rule or a validated measure of data quality. The cutoffs are round practical thresholds; no statistical identification claim follows from crossing them.

The European core contains 11 grade-A countries: Belgium, France, Germany, Italy, Netherlands (Kingdom of the), Poland, Portugal, Spain, Switzerland, Türkiye, United Kingdom. Europe follows the M49 region with Cyprus and Türkiye explicitly included in the substantive focus. Other European countries remain in the tables with their actual grades.

The four headline comparators are Australia, Japan, Singapore, United States. The United States supplies the native-classification benchmark; Australia, Japan and Singapore supply supported Asia-Pacific contrasts. This choice is based on geographic and source contrast, not estimated occupation rankings. Other grade-A countries remain available in the global inventory, rather than being treated as unsupported.

These rules constrain which evidence receives substantive emphasis; they do not assert international task equivalence. Employment-weighted statistics still differ in year, resolution and population. Countries with high occupation mass but missing compatible employment, such as Canada, retain their native occupation shares without an invented mapped employment denominator.

Grade	Geographies	Publication and employment rule
A	29	≥ 200 positive tasks; occupation mass $\geq 90\%$; two months; employment ≥ 2020 ; detail ≥ 2 digits; matched employment $\geq 80\%$
B	11	≥ 100 positive tasks; occupation mass $\geq 70\%$; same remaining conditions
C	81	May occupation evidence, but A/B conditions not all met
D	129	No May detailed occupation facet

Table 3. All 250 registry entries; grades refer to May 2026. Exact country-level components are in `outputs/evidence/geography_support.csv`.

European evidence: publication and statistical support

The table displays each component needed to read the headline European sample. Employment matching here concerns availability of a mapped score in official occupation groups. It is not the fraction of workers whose tasks have been validated or the fraction covered by an exact semantic link.

The United Kingdom uses APS employment and the ONS SOC2020 coding index. Its period label denotes the source reporting window, not a claim that all European observations share one reference date. The title-based allocation between SOC and ISCO is explicitly lexical. The machine-readable mapped-share summary reports the denominator separately for every mapping scenario.

European core	Tasks	Occupation mass (%)	Employment period	Detail	Matched (%)
Belgium	226	93.83	2025	2	99.8
France	744	98.19	2025	2	97.9
Germany	626	97.76	2025	2	99.4
Italy	420	96.78	2025	2	99.5
Netherlands (Kingdom of the)	353	95.64	2025	2	99.0
Poland	262	94.18	2025	2	98.8
Portugal	218	92.39	2025	2	99.7
Spain	528	97.46	2025	2	99.8
Switzerland	234	93.63	2025	2	95.6
Türkiye	290	94.61	2025	2	99.4
United Kingdom	680	97.67	2026-03	4	99.6

Table 4. Alphabetical presentation, without a country adoption ranking. Source: geography_support.csv. Detail is the number of classification digits; the UK retains its native SOC2020 detail.

Historical task bridge and catalogue coverage

Coverage is $C = \text{sum of positive published task indicators} / \text{catalogue task count}$. Task-share intensity is $I = \text{sum of published share fractions} / \text{catalogue task count}$. All catalogue tasks remain in the denominator, including those without published country use. The definition describes publication, not latent adoption.

Historical source texts first use the legacy-ID bridge into the current catalogue. Only source cells with no successful strict match may use exact normalized current task text as a fallback. This recovers 1,839 source cells, with 0 observed convergent collisions. Strictly matched cells are not expanded by a full union with text candidates. The strict-route and full-union diagnostics are retained separately.

Task IDs can retain changed text across catalogue vintages, and identical text can correspond to several occupation-specific tasks. These cases remain flagged. Exact text is reproducible matching evidence, not proof that task meaning is stable across time or countries.

Floors of 0.01, 0.05 and 0.10 percentage points remove small positive source cells while holding the task denominator and mapping route fixed. They are not equivalent to Anthropic's count gate because exact country conversation counts and release thresholds are not reconstructed. Missing-task scenarios and ecological automation assumptions are supplementary task diagnostics.

9. Occupational mapping changes the description

The typed ESCO-O*NET library contains 4,253 links, including 498 exact matches; the untyped alternative contains 8,627 links. Exact, exact-plus-narrow and all-typed donor means are kept separate. No exact donor yields a missing exact-route estimate. Direct ISCO group links do not become invented ESCO children.

The table compares each alternative with all typed links on pairwise common ISCO4 support. Each coefficient is a within-country occupation rank association; percentiles summarize these coefficients across the European core. Occupational sample sizes matter because the exact and untyped routes cover different universes. The complete all-May distribution is also released.

Mapping differences are material even for direct occupation shares, despite their better publication mass. A high common-support correlation can coexist with missing target groups and substantial changes for individual occupations. A route supported by broad links alone may yield a precise arithmetic mean without a precise semantic correspondence.

ISCO1 and ISCO2 are each direct means over available ISCO4 groups, avoiding a nested mean that would implicitly reweight different-sized groups. These are taxonomy weights. The scenarios and donor ranges are sensitivity evidence, not confidence intervals or guaranteed bounds on a true national occupation score. Applying the European bridge elsewhere remains an unvalidated transfer assumption.

Descriptor	Alternative	Countries	Median rho	P10 / P90	Common ISCO4
Occupation share	exact	11	0.822	0.802 / 0.842	262-262
Occupation share	exact_narrow	11	0.817	0.804 / 0.832	280-280
Occupation share	untyped	11	0.892	0.889 / 0.900	412-412
Task coverage	exact	11	0.767	0.743 / 0.812	262-262
Task coverage	exact_narrow	11	0.759	0.737 / 0.811	280-280
Task coverage	untyped	11	0.893	0.858 / 0.916	412-412

Table 5. May 2026. Source: mapping_correlation_distribution.csv. Exact_narrow means exact plus narrow links; all coefficients use pairwise common support.

Employment joins, earnings and UK mapping

For ISCO sources, select the eligible detail and retain the actual year, sex/age scope, survey population, ICLS definition and publication flags. Finer observations can be older than an available coarser series; the grade therefore checks both detail and recency. Employment weights enter only at the source's actual resolution. Unmatched or suppressed groups remain outside conditional means and are visible in the official-total denominator.

For any route, the employment-weighted descriptor is $\text{sum}(E_g \text{ times available score}_g) / \text{sum}(E_g \text{ on available-score groups})$. The matched fraction divides that denominator by the official employment total. Different routes can have different denominators. Direct-share donor means stay labelled as means; weighting them does not restore a national conversation-share allocation.

The ONS coding index contains 412 valid SOC2020 groups. The typical group links to 3 ISCO codes. Job-title counts produce lexical weights, not worker allocation probabilities. The public UK files use only the open ONS route. Restricted NFER/IBS mapping derivatives remain local.

Eurostat SES 2022 and UK ASHE provide earnings weighting sensitivities where available. Multiplying broad employee earnings means by employment from another survey does not produce a measured wage bill: units, hours, populations and years differ. Missing occupational pay is not filled with a national average. No true wage-bill exposure or national economic gain is claimed.

10. US construct comparison before any crosswalk

The US provides a direct check before an ESCO or ISCO bridge is introduced. We align native O*NET-SOC children to their six-digit SOC code, average the task metrics within each code and sum direct occupation shares across its children. Shares are summed as integer hundredths of a percentage point before conversion to fractions, preserving exact ties at the source precision. We compare the resulting May descriptors with Anthropic's published March 2026 occupation exposure file.

The associations in the table are moderate. Direct occupation shares align somewhat more closely with exposure than the two task descriptors, but none is interchangeable with the exposure score. They represent different quantities, windows and samples: monthly consumer occupation or task publication versus an exposure construction using work-related activity, global API use, capability gates, automation and time aggregation.

This is a construct-and-period comparison, not a same-input replication test. It identifies a limit of interpreting the country descriptors as exposure, without isolating which source or formula difference produces the gap. The full occupation pairs and all available US windows are released so the comparison can be examined directly.

Employment weights do not enter these rank correlations. The next section reports the separate native BLS employment join, which preserves detail without conflating the employment universe with the platform sample.

May US descriptor	Common SOC occupations	Spearman with March exposure
Published-task coverage	756	0.5869
Task-share intensity	756	0.5814
Direct occupation share	756	0.6284

Table 6. Source: us_measure_validation.csv and us_validation_pairs.csv. All 756 occupations in the published target file have finite May descriptors; unpublished mass remains distinct from actual zero use.

Native US employment and the exposure benchmark

The BLS National Employment Matrix supplies 831 detailed line items and total 2025 employment of 170,280,800 jobs. Exact native code matching supports 772 detailed groups and 92.53% of the official total. BLS-specific aggregate or unmatched codes are not split through an invented concordance. These are base-year employment counts; the 2035 projections are not used.

The Matrix counts jobs, including unincorporated self-employment, rather than individual workers (BLS, Employment Projections Data Definitions). It is not the narrower OEWS wage-and-salary universe. A frozen numeric extraction and its source hash are supplied because the automated OEWS download endpoint was unavailable. The source HTML was saved from the official BLS page, and the extraction is checked against it when the raw cache is present.

The conditional native employment-weighted task coverage is 8.977%. This is an average catalogue statistic on matched job groups. The occupation panel also reports the direct published occupation share divided by the national employment share for matching groups. This usage-to-employment ratio is a descriptive relative concentration measure, not the probability that a worker uses AI.

A separate independent public-input reconstruction of Anthropic's formula reaches Spearman 0.89465 on 756 occupations, with 6 common top-ten occupations. It fails the pre-specified 0.95 / eight-overlap gate. Simply aggregating Anthropic's already published task scores with equal weights reaches 0.87154. That downstream diagnostic embeds publisher outcomes and is not an independent reconstruction.

Unavailable original time fractions, work-share imputation, similar-task grouping, employment allocation and uncensored intermediates all obstruct exact reconstruction. Time weights are one documented missing input; the results do not identify them as the sole explanation for disagreement.

11. Theoretical exposure comparators

We compare the all-typed ISCO4 descriptors with ILO, OECD, Eloundou and Felten scores. Each correlation uses the pairwise common occupational set. The table summarizes the European core; country coefficients, all-May distributions and additional human-rated Eloundou/general-AIOE variants remain available in the data.

Positive association with some capability measures is consistent with usage concentrating in tasks that language models can assist. It does not show that a share is a time-exposure fraction or validate a causal employment channel. The direct occupation and task descriptors can differ in their alignment because they combine different information about breadth and concentration.

The OECD input is already the published reversed normalized exposure measure, with higher values denoting greater exposure. Its negative associations are retained. Flipping that sign to obtain a familiar result would alter the source construct. These data alone cannot distinguish the contributions of capability definitions, classification mapping and selected platform activity to the disagreement.

The table contains aggregate correlations only. Restricted source score vectors and restricted mapping-dependent derivatives are excluded from the public archive; source-specific attribution and redistribution rules remain documented.

Comparator	Descriptor	Countries	Median rho	P10 / P90	Common occupations
ILO	Occupation share	11	0.630	0.617 / 0.654	406-406
ILO	Task coverage	11	0.580	0.527 / 0.605	406-406
OECD	Occupation share	11	-0.136	-0.180 / -0.120	411-411
OECD	Task coverage	11	-0.106	-0.197 / -0.051	411-411
Eloundou GPT-4 beta	Occupation share	11	0.666	0.642 / 0.705	412-412
Eloundou GPT-4 beta	Task coverage	11	0.608	0.559 / 0.655	412-412
Felten language	Occupation share	11	0.701	0.675 / 0.746	411-411
Felten language	Task coverage	11	0.634	0.568 / 0.691	411-411

Table 7. May 2026, common ISCO4 groups, all-typed mapping. Source: comparator_distribution.csv. Country counts are not independent observations of a national workforce.

12. PIAAC: agreement across occupations

PIAAC supplies independent survey evidence on the frequency of broad work activities. The frozen bridge connects survey items to seven O*NET activity domains. The primary descriptive specification compares the survey share reporting an activity at least weekly with its share of catalogue tasks, using all typed links at ISCO2. Each correlation is calculated across occupation groups within one survey geography and one activity domain.

The headline statistic is therefore across occupations, with the number of occupation groups shown in the table. It is not the small correlation across five to seven domains inside one occupation. We retain every finite coefficient, including coefficients flagged for weak alignment, rather than selecting rows according to their sign or review status. Country means are unweighted means of coefficients, not pooled person-level correlations.

The strength of agreement varies across activities. Analytical and documentation profiles align more closely than communication or numerical-processing profiles in this specification. These differences reveal limits in transferring a US task catalogue to survey activity descriptions. They do not establish task-time shares or country AI adoption.

Each public activity cell requires at least 30 valid respondents; the across-occupation calculation also requires the protocol's minimum common occupational support. Survey weights enter the activity shares, but no replicate-weight design interval is claimed for these rank summaries. England and other restricted survey scopes retain their own labels. Daily-frequency outcomes and alternative O*NET profiles are released separately.

Activity domain	Survey geographies	Mean rho	Median rho	Occupations per correlation
analysis	16	0.733	0.740	26-35
communication	16	0.373	0.390	19-31
computer work	16	0.479	0.470	19-31
documentation	16	0.568	0.587	26-35
information	16	0.472	0.464	18-31
measurement	16	0.390	0.428	26-35
numerical processing	16	0.290	0.283	26-35

Table 8. ISCO2; at-least-weekly survey share versus catalogue task share; all-typed route. Source: `piaac_across_occupation_summary.csv`. Occupation counts, not respondent totals, determine each rank comparison.

Detailed occupational profiles: France and Spain

Four-digit results provide a useful check on whether broad-group agreement survives more detailed occupational comparison. France and Spain are shown separately using exactly the same weekly catalogue-share specification as the two-digit table. The underlying number of common occupations varies by domain and country; those denominators are an essential part of the result.

Fine-detail agreement is uneven, including near-zero or negative associations for some French domains. Averaging these results into one universal validation coefficient would conceal the variation. Sparse occupational support also limits how broadly the detailed findings can be generalized.

The within-occupation profile correlations remain supplementary diagnostics. Their small number of domains makes a count below an arbitrary correlation threshold unsuitable as the principal validation headline. The broader across-occupation results reported here still constitute a partial content check, not a validated time-use model.

Domain	France rho	France occupations	Spain rho	Spain occupations
analysis	0.630	33	0.500	32
communication	0.410	27	0.275	14
computer work	0.358	27	0.433	14
documentation	0.500	33	0.502	32
information	-0.073	27	0.291	14
measurement	-0.044	33	0.361	32
numerical processing	-0.022	33	0.233	32

Table 9. ISCO4; all-typed mapping; at-least-weekly survey share versus catalogue task share. Source: `piaac_detailed_country_results.csv`. No filtering on coefficient sign or review flag.

13. Time-weight evidence and research protocol

A frozen local-model task-hour experiment fails basic plausibility. One task received two quadrillion weekly hours, and twelve occupation totals exceeded 168 hours. Structural validity, exact task IDs and normalized shares cannot make those quantities credible. The original values and their arithmetic propagation remain available with failure flags as an audit record.

The failed weights and their exposure correlations are excluded from substantive conclusions. In particular, the experiment supplies no evidence that time weights are unimportant or that a stronger model would resolve the reconstruction gap. We do not clip an extreme response, selectively regenerate occupations or select weights by agreement with the target.

The prospective protocol in docs/time_weight_rerun_protocol.md specifies three blinded independent draws, an explicit 40-hour analytical allocation including uncovered activity, exact task IDs, nonnegative finite entries, sum constraints and whole-draw rejection rules. The mean is defined only for occupations with three valid draws. Missingness and between-draw variation must be reported.

Before a confirmatory run, the endpoint and settings, independent human time-use evidence, validation sample and success criteria must be frozen and publicly registered. The current document is a local prospective design, not an executed or externally preregistered study. Model capacity alone is not external validation. PIAAC frequency responses and O*NET importance ratings cannot substitute for observed time allocations.

Anthropic exposure algebra and identification limits

The published method combines work-related consumer task counts C_{work} and global API counts A as $W = C_{\text{work}} + A$. For $W > 0$ the automation factor is $\alpha = 0.5 + 0.5 \times (C_{\text{work}} \times \text{automation share} + A) / W$. Task exposure applies the $W \geq 100$ and capability $\beta \geq 0.5$ gates, then aggregates with occupation-specific normalized time fractions. The public proxy defines the zero- W case explicitly.

The fixed independent proxy has Spearman 0.89465 and fails the 0.95 / eight-common-top-ten gate. Published-task-score equal aggregation reaches 0.87154, also below that threshold, but uses task outcomes already constructed by the publisher. It checks a downstream aggregation step, not the independent input pipeline.

Missing original time fractions are a verified obstacle, alongside work-share imputation, similar-task groups, employment allocations and unrounded or uncensored intermediates. Reproducing one missing input would not establish complete reconstruction. The failed model experiment neither resolves the obstacle nor ranks these missing inputs by importance.

The direct occupation panel is deliberately defined from published source fields. Its validity as a published-mass descriptor does not depend on passing an exposure-equivalence gate, but it must not be described as that exposure measure. The US comparison makes the construct difference observable.

14. Interpretation and implications

The combined findings support a hierarchy of uses. Country usage share and the publisher's per-capita index describe the relative scale of platform activity. Direct occupation shares describe published occupational concentration with substantially better retained mass than task rows. Task coverage describes breadth in a fixed catalogue and is strongly related to country usage volume. These roles should remain distinct in empirical applications.

For Europe, explicit classification routes are necessary but do not establish local task equivalence. An exact semantic link, an equal taxonomy mean and an employment weight answer different questions. Reporting support, alternative routes and common occupational samples prevents arithmetic precision from being mistaken for measurement accuracy. The US native route illustrates what can be checked before a European bridge is introduced.

The UK threshold diagnostic makes the breadth issue concrete: on the same ONS mapping and APS employment support, May task coverage is 6.0209% at baseline and 1.4994% after a 0.10 percentage-point source-share floor. This is sensitivity to small published cells, not evidence of a change in adoption.

The data can support occupation-focused descriptive work and help select questions for worker or firm studies. Estimating employment effects requires an outcome design and assumptions about exposure, timing and confounding. Estimating work-time automation requires independently validated time use and reliable work-related classifications. Neither task floors nor employment reweighting recovers these missing ingredients.

The global inventory remains useful precisely because it does not force sparse, old or incompatible evidence into a common headline. Grades and the European focus constrain interpretation while retaining machine-readable observations for all supported geographies. The study's scientific claim is the documented measurement evidence, including its failures and unresolved inputs.

Files, audit records and release practice

Begin with `outputs/evidence/summary.json`, `geography_support.csv` and `facet_comparison.csv`. `country_usage.csv` contains the overall source fields; `direct_occupation_source.csv` retains published rows. `onet_occupation_panel.csv`, `esco_occupation_panel.csv` and `isco_occupation_panel.csv` distinguish native shares from donor means. `us_native_soc_panel.csv` supplies BLS exact-code joins and usage-to-employment ratios.

`volume_regression.csv`, `us_measure_validation.csv`, `mapping_correlation_distribution.csv`, `comparator_distribution.csv` and the two PIAAC summary files generate the manuscript tables. Each family retains underlying comparison rows and sample counts. `dataset_catalogue.csv` and `data_dictionary.csv` describe keys, units and public bytes. The task panels remain available as secondary diagnostics.

The build preserves source hashes and checks unique keys, range restrictions, missingness and aggregation denominators. Numerical reproduction, public-source rights and PDF layout checks are distinct from scientific validation. Their receipts identify exactly which files were checked and must be refreshed when those files change.

Source vintages, source commits and file hashes identify the data and transformations represented in the research. The website provides the paper, technical note and licensed dataset subset. Publication status does not establish exposure equivalence or causal validity.

15. Reproducibility, attribution and access

All numerical manuscript tables are read by the report builder from computed CSV outputs. The source manifest records URLs, retrieval times and hashes; the data catalogue records keys, units, row counts and public-file hashes. Code separates source publication flags, taxonomy aggregation, official employment joins and descriptive validation.

The public package contains the licensed subset of data and code. Restricted mapping vectors, restricted comparator vectors and respondent-level PIAAC records are not redistributed. The frozen BLS base-year extraction includes the source HTML hash and extraction description. Failed model weights are retained for audit and deterministic arithmetic reproduction, with no inference required to rebuild the analyses.

The direct occupation evidence and reporting grades are documented in the paper, technical note and evidence tables. Supplementary task-coverage diagnostics describe published catalogue breadth; they do not constitute a validated country exposure ranking.

Source identifiers, manifests and file hashes preserve the evidence behind the calculations. Fatih Kansoy is the author of this independent research manuscript. Computational reproducibility is separate from scholarly validation.

References

Handa, K., et al. (2025). Which Economic Tasks are Performed with AI? Evidence from Millions of Claude Conversations. arXiv:2503.04761.

<https://arxiv.org/abs/2503.04761>

Appel, R., McCrory, P., Tamkin, A., McCain, M., Neylon, T., and Stern, M. (2025). Anthropic Economic Index report: Uneven geographic and enterprise AI adoption. arXiv:2511.15080.

<https://arxiv.org/abs/2511.15080>

Anthropic (2026). Anthropic Economic Index report: Cadences. 26 June. Country data documentation and pinned EconomicIndex release accompany the report.

<https://www.anthropic.com/research/economic-index-june-2026-report>

Anthropic (2026). EconomicIndex dataset, revision 2ea58ff75e4247d26810c37f10c179edc2466cac. Source directories and file hashes are in MANIFEST.csv.

<https://huggingface.co/datasets/Anthropic/EconomicIndex/tree/2ea58ff75e4247d26810c37f10c179edc2466cac>

Anthropic (2026). Labor market impacts of AI: A new measure and early evidence. March report and methodology appendix.

<https://www.anthropic.com/research/labor-market-impacts>

References (continued)

Chatterji, A., Cunningham, T., Deming, D. J., Hitzig, Z., Ong, C., Shan, C. Y., and Wadman, K. (2025). How People Use ChatGPT. NBER Working Paper 34255.

<https://www.nber.org/papers/w34255>

Tomlinson, K., Jaffe, S., Wang, W., Counts, S., and Suri, S. (2025). Working with AI: Measuring the Applicability of Generative AI to Occupations. arXiv:2507.07935.

<https://arxiv.org/abs/2507.07935>

Brynjolfsson, E., Li, D., and Raymond, L. (2025). Generative AI at Work. Quarterly Journal of Economics, 140(2), 889-942.

<https://doi.org/10.1093/qje/qjae044>

Humlum, A., and Vestergaard, E. (2025). Large Language Models, Small Labor Market Effects. NBER Working Paper 33777, original title; the current landing page uses Still Waters, Rapid Currents: Early Labor Market Transformation under Generative AI.

<https://www.nber.org/papers/w33777>

Bick, A., Blandin, A., and Deming, D. J. (2024; revised 2025). The Rapid Adoption of Generative AI. NBER Working Paper 32966.

<https://www.nber.org/papers/w32966>

References (continued)

Fan, R. Y., and Nguyen, H. M. (2026). Aggregate Gains from AI and Their Distribution: Global Evidence from Usage Data. IMF Working Paper 2026/147.

<https://www.imf.org/en/publications/wp/issues/2026/07/10/aggregate-gains-from-ai-and-their-distribution-global-evidence-from-usage-data-577586>

Steele, J. L., and Cruz, I. (2026). Helping People Choose Careers in the Age of AI. arXiv:2607.15506.

<https://arxiv.org/abs/2607.15506>

Eloundou, T., Manning, S., Mishkin, P., and Rock, D. (2023). GPTs are GPTs: An Early Look at the Labor Market Impact Potential of Large Language Models. arXiv:2303.10130.

<https://arxiv.org/abs/2303.10130>

Felten, E., Raj, M., and Seamans, R. (2021). Occupational, industry, and geographic exposure to artificial intelligence: A novel dataset and its potential uses. Strategic Management Journal.

doi:10.1002/smj.3286.

<https://doi.org/10.1002/smj.3286>

Felten, E., Raj, M., and Seamans, R. (2023). How will Language Modelers like ChatGPT Affect Occupations and Industries? arXiv:2303.01157.

<https://arxiv.org/abs/2303.01157>

References (continued)

Centre for British Progress (2026). AI and the UK labour market: the evidence so far. April report and downloadable chart data.

<https://britishprogress.org/reports/ai-and-the-uk-labour-market-the-evidence-so-far>

Bank Underground (2026). Canaries in the column: AI exposure and the UK's hiring slowdown. 6 August. Blog post; views are those of its authors.

<https://bankunderground.co.uk/2026/08/06/canaries-in-the-column-ai-exposure-and-the-uks-hiring-slowdown/>

European Commission. Crosswalk between ESCO and O*NET. Typed semantic links and release documentation.

<https://esco.ec.europa.eu/en/about-esco/data-science-and-esco/crosswalk-between-esco-and-onet>

O*NET Resource Center. O*NET Database 30.2 and archived releases. Pinned occupational task catalogue.

https://www.onetcenter.org/db_releases.html

Office for National Statistics. SOC2020 Volume 2: coding rules and conventions; occupation coding index. APS employment is obtained through Nomis.

<https://www.ons.gov.uk/methodology/classificationsandstandards/standardoccupationalclassificationsoc/soc2020/soc2020volume2codingrulesandconventions>

References (continued)

US Bureau of Labor Statistics (2026). Table 1.2: Occupational projections, 2025-35, and worker characteristics, 2025. National Employment Matrix base-year employment only.

<https://www.bls.gov/emp/tables/occupational-projections-and-characteristics.htm>

US Bureau of Labor Statistics. Employment Projections Data Definitions. Matrix employment counts jobs, including self-employment, rather than individual workers.

<https://www.bls.gov/emp/documentation/definitions.htm>

International Labour Organization (2025). Generative AI and Jobs: A Refined Global Index of Occupational Exposure. ILO Working Paper 140.

https://www.ilo.org/sites/default/files/2025-05/WP140_web.pdf

OECD. The OECD AI exposure measure. Published occupation scores and methodological report.

https://www.oecd.org/en/publications/the-oecd-ai-exposure-measure_f3da0f0a-en.html

OECD. Survey of Adult Skills (PIAAC), Cycle 2 database, questionnaires and country public-use-file documentation.

<https://www.oecd.org/en/data/datasets/piaac-2nd-cycle-database.html>

References (continued)

Eurostat. EU Labour Force Survey employment by occupation (lfsa_egai2d) and Structure of Earnings Survey 2022. Source population and quality flags retained.

https://ec.europa.eu/eurostat/databrowser/view/lfsa_egai2d/default/table?lang=en

ILOSTAT. Annual employment by sex and occupation, bulk dissemination facility. ISCO-08 selection and country source identifiers retained.

<https://rplumber.ilo.org/data/indicator/>