



EW-AiRM™
BY HUMAN-AI.INSTITUTE

FORMAL SUBMISSION

ATUS Artificial Intelligence Questions by the U.S. Department of Labor

Public comment to the Bureau of Labor Statistics, U.S. Department of Labor, on the proposed information collection *American Time Use Survey (ATUS) Artificial Intelligence (AI) Questions*, applying EW-AiRM™ (Enterprise-Wide AI Risk Management) to the four desired focus areas of the Paperwork Reduction Act pre-clearance consultation

"Governance that looks right is not the same as governance that works."

EW-AIRM™ GOVERNANCE MAXIM — THE SAME DISCIPLINE APPLIES TO MEASUREMENT

SUBMITTED BY

Prof. Markus Krebsz
Creator & Author
of EW-AiRM™
Founding Director,
Human-Ai.Institute

DATE

3 September 2026
Comments close
8 September 2026

SUBMITTED TO

Bureau of Labor
Statistics,
U.S. Department of Labor
91 FR 42775
OMB 1220-NEW

ONLINE

www.ewairm.com
www.human-ai.institute

© 2026 Prof. Markus Krebsz · Human-Ai.Institute · De-Risking Solutions Ltd. All rights reserved. EW-AiRM™, HAIPECR™ and the associated logos are protected marks.

The author consents to this comment being treated as a public record and published in full, including in the request for Office of Management and Budget approval. The views expressed are the author's own. This comment contains no confidential information.

Prof. Markus Krebsz

Creator & Author, EW-AiRM™ (Enterprise-Wide AI Risk Management)
Founding Director, Human-Ai.Institute · UNECE WP.6 AI Project Lead
contact@human-ai.institute · www.ewairm.com



Ms Erin Good
BLS Clearance Officer
Division of Management Systems
Bureau of Labor Statistics, U.S. Department of Labor
Washington, DC, United States of America
By email: BLS_PRA_Public@bls.gov

3 September 2026

Re: Proposed Information Collection; ATUS Artificial Intelligence (AI) Questions, 91 FR 42775 (10 July 2026), FR Document 2026-13928, OMB Control Number 1220-NEW

Dear Ms Good,

Thank you for the opportunity to comment on the Bureau of Labor Statistics' proposed new information collection, the American Time Use Survey (ATUS) Artificial Intelligence (AI) Questions. I submit this response in my personal professional capacity as Founding Director of the Human-Ai.Institute and as creator and author of **EW-AiRM™ (Enterprise-Wide AI Risk Management)**, a three-layer AI risk governance framework published in full by Wiley Finance on 12 November 2026 and listed in the OECD.AI Catalogue of Tools and Metrics for Trustworthy AI.

My headline position is one of **strong support**. The Bureau is right that no existing federal survey links individuals' AI use to a continuous, activity-level record of their day, and the ATUS time diary is the correct instrument to close that gap. This collection would do something no vendor survey or platform telemetry release can do: provide a nationally representative, methodologically transparent baseline of who uses AI, for which tasks, and with what displacement of time, against which organisations, researchers, insurers, and standard-setters can calibrate their own measurements.

The attached comment responds to each of the Bureau's four desired focus areas and offers ten specific recommendations, the most consequential of which concern question design: distinguishing generative, embedded, and agentic AI use; capturing use that respondents do not recognise as AI; measuring whether AI outputs are verified before use; and identifying unsanctioned workplace AI use, often called shadow AI, for which no national estimate currently exists.

I consent to this comment being treated as a public record and published in full. It contains no personal or confidential information beyond my professional details. I would be pleased to discuss any aspect of this submission with Bureau staff, including by video-conference from the United Kingdom, and to share the underlying EW-AiRM™ measurement constructs (i.e. KXIs) should that be useful to the question-design and cognitive-testing phases.

Yours sincerely,

Prof. Markus Krebsz

UN ECE	AI Project Lead (Embedded AI Systems)
The Human AI Institute	Founding Director
Enterprise-wide AI Risk Management (EW-AiRM)	Creator & Author
Stirling University (UK)	Honorary Professor
Woxsen University (India)	Clinical Professor of Practice

unece.org/trade/wp6/digital-regulation-goods-artificial-intelligence
unece.org/trade/documents/2025/03/informal-documents/markus-krebsz-wp6-project-leader-ai
www.Human-Ai.Institute · www.ewairm.com · www.stir.ac.uk/expert/in/artificial-intelligence · www.krebsz.net

contact@human-ai.institute · krebsz@web.de · +44 (0) 79 85 065 045

About the Submitter and About EW-AiRM™

1.1 The submitter

Prof. Markus Krebsz is the creator and author of EW-AiRM™ (Enterprise -Wide AI Risk Management) and the Founding Director of the **Human -Ai.Institute** and of De-Risking Solutions Ltd. He served as AI Project Lead for the United Nations Economic Commission for Europe (UNECE) Working Party 6 on Regulatory Cooperation and Standardization Policies , where he led the development of the Common Regulatory Arrangement on products with embedded artificial intelligence (ECE/TRADE /486), adopted in 2024 and available to the 56 UNECE member states , including the United States . He is a member of the European AI Office (DG CNECT) expert group , a plenary participant in the EU General-Purpose AI Code of Practice process and a member of all four of its working groups, a StandICT .eu Standards Maker contributing to IEEE P7007.1 (Ontological Standard for Addressing Risks in AI Systems), a ForHumanity Certified Auditor , an Honorary Professor at the University of Stirling (UK) and a Clinical Professor of Practice at Woxsen University (India). He has served since 2017 as an independent non-executive director on the Supervisory Council of a European licensed bank.

The Human- Ai.Institute is an inaugural member of the United Nations University Global AI Network and a founding member of UNESCO 's Internet for Trust network . Further information is available at www.ewairm.com , www.human-ai.institute and www.krebsz.net.

This comment is made because household -level statistics on AI use are a direct input to enterprise AI risk management . The organisations that adopt frameworks such as EW-AiRM™ currently have no reliable national reference point for how, where, and how much AI is actually used by the people who work for them. The proposed ATUS module would create one.

1.2 What EW-AiRM™ is

EW-AiRM™ is a three-layer framework for managing AI risks in any organisation , designed to be proportionate across three tiers of organisational scale (Core, Standard, and Full). Its **Strategic layer** comprises six pillars, from strategic alignment and necessity assessment through governance and accountability to through-the- lifecycle monitoring . Its **Operational layer** is built on the MIT AI Risk Repository (MIT FutureTech), which categorises more than 1,700 AI risks across seven domains and twenty -four subdomains . Its **Resilience layer** addresses eight categories of low-probability, high-impact AI Black Swan events, including multi-agent emergence and systemic foundation model concentration . Across all three layers runs **HAiPECR™**, a seven - dimension ethical overlay originating in 2023 as the **Human-Ai Paradigm for Ethics, Conduct and Risk**, mapped to the ten core principles of the UNESCO Recommendation on the Ethics of Artificial Intelligence (2021) and listed on the OECD AI Policy Observatory since April 2023. Measurement uses **Key X Indicators (KXIs)**, the umbrella for key performance, key risk and key control indicators.

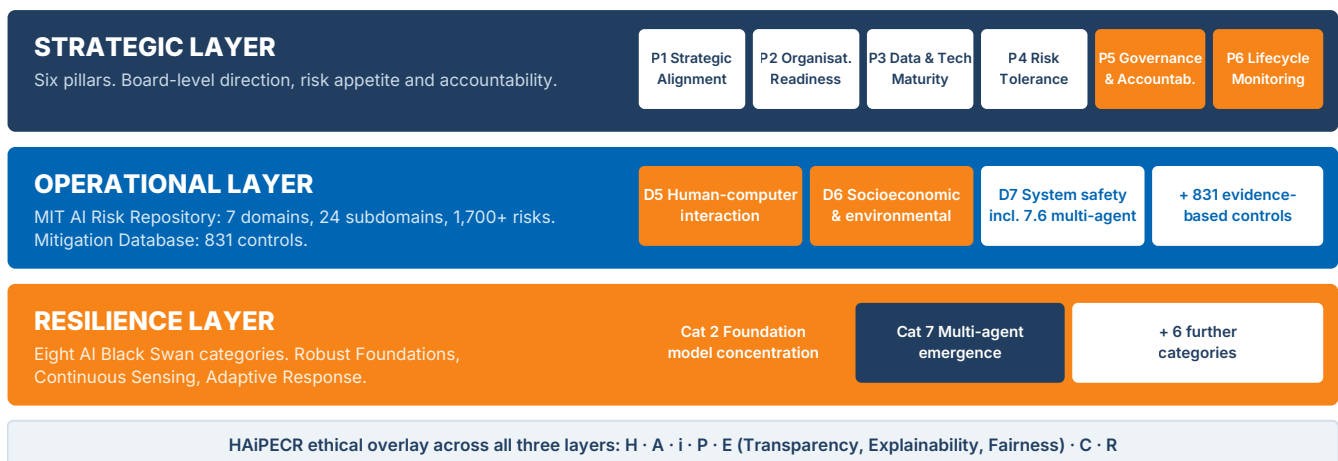


Figure 1. The three-layer architecture of EW-AiRM™. Elements highlighted in orange and navy are those most directly engaged by this consultation: Pillars 5 and 6, operational domains D5 and D6, and Black Swan Categories 2 (foundation model concentration, see Section 6.1) and 7 (multi-agent emergence), which falls within the proposed two-year fielding window. © 2026 Prof. Markus Krebsz. (Displaying only a selection of EW-AiRM™ framework building blocks that are relevant to this consultation.)

1.3 Lineage and international standing

EW-AiRM™ is the culmination of more than a decade of governance instrument development by the author. It builds on the Global Conduct Risk Paradigm (GCRP, 2015); the Universal Conduct Risk Paradigm (UCRP), presented to the UNECE Group of Experts on Risk Management in Regulatory Systems in February 2017 and published by UNECE; HAiPECR (2023), listed on the OECD AI Policy Observatory; and the UNECE Common Regulatory Arrangement on products with embedded AI (ECE/TRADE/486, 2024), for which the author served as project lead and which is the regulatory-side counterpart to EW-AiRM™ at the enterprise level. EW-AiRM™ is now also listed in the OECD.AI Catalogue of Tools and Metrics for Trustworthy AI and is designed for direct interoperability with ISO/IEC 42001, the NIST AI Risk Management Framework, the EU AI Act and the OECD AI Principles, to which the United States is a founding adherent. This lineage matters for the Bureau's purposes: the **framework has been stress-tested through intergovernmental negotiation and regulated-industry practice**, and its measurement constructs are designed to interoperate with, rather than compete against, official statistics.



Figure 2. Lineage of EW-AiRM™: from conduct-risk paradigms to an enterprise-wide AI risk management framework with regulatory-side and enterprise-side instruments. © 2026 Prof. Markus Krebsz.

Executive Summary and Recommendations

In one paragraph. The Bureau of Labor Statistics is right, and in this author's assessment understates its own case: linking AI use to a continuous 24-hour activity record is not merely a useful supplement to existing technology statistics, it is the only survey design capable of measuring what actually matters about AI adoption, which is task substitution and time reallocation rather than headline adoption incidence. Enterprise-side surveys tell us whether organisations say they use AI; the ATUS module would tell us what people actually do with it, minute-by-minute, across work, education, household production, and leisure. Those are different facts, and risk management needs both. This submission **strongly supports approval** of the collection and concentrates its recommendations on question design, because the analytical value of the entire two-year collection will be determined by a small number of definitional choices made before January 2027. The single greatest threat to the collection's practical utility is not burden, at an estimated two minutes per response one of the lightest modules the Bureau will field, but **construct drift**: AI is disappearing into ordinary software faster than respondents can recognise it, and a module that measures only self-identified, deliberate use of a chatbot will systematically undercount AI exposure, more severely in the second year of fielding than the first.

Ten recommendations

Each recommendation is tagged to the Bureau's desired focus areas: [A] necessity and practical utility · [B] accuracy of the burden estimate · [C] quality, utility, and clarity · [D] minimising burden.

R1

Proceed with the collection. [A] The design is necessary, has practical utility well beyond the Bureau's own publications, and fills a gap no other federal instrument can fill. The January 2027 start date is well judged: it captures the adoption curve while it is still steep.

R2

Define "AI tool" operationally, with a three-way taxonomy. [C] Distinguish (i) general-purpose generative AI used deliberately, (ii) AI features embedded within other software, and (iii) agentic AI acting with a degree of autonomy on the respondent's behalf. These carry materially different risk profiles and different labor-market implications, and agentic use will grow substantially within the fielding window.

R3

Capture unrecognised use. [C] Include a "don't know" pathway and a short awareness probe so that respondents who cannot tell whether a tool involved AI are measured rather than silently misclassified as non-users. Document the awareness rate in the technical notes so researchers can bound the undercount.

R4

Add one verification-behaviour item. [C] Ask whether the respondent reviewed, corrected, or checked AI outputs before using them. This single item transforms the module's value for productivity, quality-of-work, and risk research, because time spent verifying is the hidden denominator of every AI productivity claim.

R5

Distinguish sanctioned from unsanctioned workplace use. [C] Where AI was used for work, ask whether the tool was provided or approved by the employer or was the respondent's own. This would produce the first nationally representative estimate of so-called shadow AI, a quantity every enterprise risk framework, including EW-AiRM™, currently has to assume rather than know.

R6

Harmonise the core screener. [C] Align the wording of the basic adoption item with the Census Bureau's firm-side AI measures and with emerging international household-measurement work, so household and business series can be triangulated and U.S. data compared internationally.

R7

Freeze wording, refresh examples. **C** Hold question wording constant across the full two-year window to preserve trend integrity, and manage the fast-moving tool landscape by updating illustrative examples only, with every change documented in the technical documentation.

R8

Validate the two-minute estimate empirically. **B** The estimate is plausible for a gated design but sensitive to definitional deliberation by respondents. Confirm it through cognitive testing before fielding and publish observed administration times in the first technical documentation release.

R9

Gate and branch. **D** Use a single screener with the full module administered only to AI users, placed after the diary so that task probes can anchor to activities the respondent has already reported rather than asking anything twice.

R10

Publish a full crosswalk with the microdata. **A** **C** Release the public-use microdata with an explicit crosswalk between the AI items and ATUS activity lexicon codes, and document any automated or AI-assisted processing applied to responses, as an act of the same transparency the data will enable in others.

Section 6 adds four forward-looking observations that sit beyond the four statutory criteria but within the design horizon of the collection. They are offered as observations rather than recommendations, so that the ten recommendations above can be weighed cleanly against the criteria the Paperwork Reduction Act prescribes.

Why Activity-Linked AI Data Matters for Enterprise AI Risk Management

The Bureau frames the collection's utility in terms of productivity, job tasks, skill requirements, and time allocation. All correct. This submission adds a further category of practical utility that the notice does not claim for itself but that materially strengthens the case for approval: **the collection is a missing input to organisational AI governance.**

3.1 Enterprises are currently governing blind

Under EW-AiRM™, as under the NIST AI Risk Management Framework and ISO/IEC 42001, an organisation's first governance obligation is to know what AI is actually in use. EW-AiRM™ makes this concrete through its Five Non-Negotiables: every deployed AI system must have a named individual owner, a pre-deployment assessment, a tested human override capability, an explicit acceptance of residual risk, and an incident reporting pathway. Each of these presupposes that the organisation can see its own AI use. In practice, it frequently cannot, because a growing share of workplace AI use happens through personal accounts and consumer tools that never pass through procurement. Industry surveys repeatedly suggest this unsanctioned use is widespread, but every such survey is a convenience sample commissioned by a vendor with a product to sell. There is **no nationally representative estimate of shadow AI in the United States**. Recommendation R5 would create one, at the cost of a single question.

EW - AIRM™ RISK LOCATION

Shadow AI is a direct threat to Non-Negotiable NN1 (named ownership of every deployed system) and NN2 (pre-deployment assessment): a tool the organisation cannot see cannot be owned, assessed, or overridden. In the Operational layer it cuts across the MIT AI Risk Repository's human-computer interaction (D5) and security domains. National prevalence data would let organisations benchmark their internal discovery efforts against a trustworthy baseline for the first time.

3.2 Time is the honest metric

Adoption statistics answer "whether"; time-use statistics answer "how much and instead of what". For risk purposes the second question dominates. A workforce in which 60 percent of people have tried a chatbot once is a very different exposure from a workforce in which 15 percent of people route two hours of daily work through one. The ATUS diary linkage is the only federal design that can separate these worlds. It is also the only design that can measure the verification time discussed in Recommendation R4: if AI saves forty minutes of drafting but adds twenty-five minutes of checking, that fact is invisible to every adoption survey and fully visible to a time diary.

"Hallucination is a property of the architecture, not a failure of it."

EW-AiRM™ GOVERNANCE MAXIM · KREBSZ, ENTERPRISE-WIDE AI RISK MANAGEMENT (WILEY FINANCE, 2026). WHICH IS PRECISELY WHY VERIFICATION BEHAVIOUR, NOT ADOPTION, DETERMINES WHETHER AI USE CREATES OR DESTROYS VALUE.

3.3 From national statistics to enterprise instrumentation

EW-AiRM™'s monitoring pillar operates through **Key X Indicators (KXIs)**, an umbrella construct spanning key performance, key risk, and key control indicators. Several standard KXIs, including workforce AI adoption rate, task coverage, and oversight-mode distribution, currently lack any external reference point: an organisation can

compute its own figure but cannot say whether that figure is high, low, or typical. The proposed ATUS module would supply the missing denominator. Figure 3 shows the transmission path.



Figure 3. Transmission path from the proposed ATUS collection to enterprise AI governance under EW-AiRM™. National statistics supply the external denominators that enterprise Key X Indicators (KXIs) currently lack. © 2026 Prof. Markus Krebsz.

3.4 A two-year window that will span a structural shift

The two-year fielding window, January 2027 through December 2028, will almost certainly span the mainstreaming of agentic AI: systems that do not merely respond to prompts but plan and execute multi-step tasks, including transactions, on a person’s behalf. In EW-AiRM™’s Resilience layer, multi-agent emergence is one of the eight AI Black Swan categories precisely because interacting autonomous systems produce behaviour that no single deployer intends or anticipates. A module designed in 2026 around the mental model of “chatting with an AI” will misdescribe 2028. This is the strongest argument for Recommendation R2’s three-way taxonomy: it future-proofs the instrument without changing its wording mid-series.

Responses to the Four Desired Focus Areas

4.1 Necessity and practical utility of the collection [Focus area A]

The collection is necessary and its practical utility is high. The Bureau's statutory mission under 29 U.S.C. 2 extends to emerging economic trends, and it is difficult to name a trend more consequential for time allocation, job tasks, and skill requirements than AI diffusion. The Bureau's claim that no other national survey provides activity-level detail linking AI use to a continuous 24-hour record is accurate: firm-side instruments measure organisational adoption, and household technology supplements measure access and incidence, but neither can observe substitution within a person's day.

Beyond the Bureau's own tabulations, four external constituencies have concrete, current demand for exactly this data: **(i) enterprise risk and governance functions**, which need national baselines to calibrate internal indicators (Section 3.3); **(ii) insurers and actuaries**, who are beginning to underwrite AI-related operational and liability exposures with no population-level usage data; **(iii) standards bodies and regulators**, domestic and international, whose proportionality judgements about AI rules currently rest on vendor-published figures; and **(iv) the research community**, for whom the diary linkage enables task-exposure analysis that occupational classifications alone cannot support. The production of public-use microdata files multiplies this utility; Recommendation R10 addresses how to maximise it.

4.2 Accuracy of the burden estimate [Focus area B]

The estimate of 7,672 annual respondents at an average of two minutes, for 256 annual burden hours, is **plausible but sensitive to design choices** in two respects.

- **Definitional deliberation is burden.** The dominant driver of response time for this module will not be the number of items but the time respondents spend deciding whether something counts as AI. "Did I use AI yesterday?" is, in 2027, a genuinely hard question: autocomplete, navigation, photo search, spam filtering, and voice assistants all involve AI that respondents may or may not recognise as such. A crisp operational definition with concrete examples (Recommendation R2) is therefore simultaneously a clarity improvement and a burden reduction. Ambiguity taxes every respondent; definition taxes none.
- **The average conceals a split.** If the module is gated (Recommendation R9), non-users will complete it in well under a minute while users receiving task probes may exceed two minutes, particularly where probes iterate over diary activities. The average may hold while the user-side experience does not. The Bureau should validate the estimate through cognitive and timing tests before fielding and publish observed administration times in the first technical documentation release (Recommendation R8), which would also set a useful precedent for AI-related collections government-wide.

For the record: 256 burden hours is an exceptionally modest cost for the first nationally representative, activity-linked AI usage series in the world. On any reasonable weighing, the balance falls decisively in favour of collection.

4.3 Quality, utility, and clarity of the information to be collected [Focus area C]

This is where the collection will be won or lost, and where EW-AiRM™ has the most to offer. Five design issues follow.

4.3.1 A three-way taxonomy of AI use (R2)

The module should distinguish, at minimum: **deliberate generative use** (the respondent knowingly used a general-purpose AI tool such as a chatbot or image generator); **embedded use** (the respondent used software

in which AI features are integrated, such as drafting suggestions, summarisation, or search assistance inside productivity tools); and **agentic use** (an AI system carried out a multi-step task on the respondent's behalf, such as scheduling, purchasing, or filing). The three modes differ in who initiates the activity, how visible the AI is, and where accountability sits, which is why EW-AiRM™'s Operational layer treats them under distinct risk domains. They also differ in labor-market meaning: embedded use augments a task the human still performs; agentic use replaces the performance itself. A module that pools them will produce a single series that means less each year.

4.3.2 Measuring what respondents cannot see (R3)

Survey measurement of AI use has an unusual property: the measurement error is trending. As AI is absorbed invisibly into ordinary software, the gap between actual and self-recognised use widens, which means an instrument held constant will show spurious plateaus. The remedy is not to chase the technology with rewording, which destroys the trend, but to **measure the awareness gap itself**: retain a "don't know" response pathway, include one short probe about uncertainty, and publish the resulting awareness rates so analysts can bound the undercount. This is the survey-methodology equivalent of EW-AiRM™'s through-the-lifecycle monitoring pillar: instrument the drift rather than pretend it away.

EW - A I R M ™ R I S K L O C A T I O N

Pillar 6 (Through-the-Lifecycle Monitoring and Adaptability) requires that measurement instruments be monitored for validity drift just as models are monitored for performance drift. The same principle applies to a national statistical instrument observing a moving technological target: the drift must be measured, documented, and disclosed, not absorbed silently into the series.

4.3.3 Verification behaviour (R4)

A single item asking whether the respondent checked, corrected, or edited AI outputs before using them would be, in this author's assessment, the highest-value question per second of burden in the entire module. It speaks directly to productivity measurement (verification time offsets generation time), to quality and safety research (unverified use in consequential tasks is a leading indicator of AI-related incidents), and to skills policy (verification is itself the emerging skill). No national survey anywhere currently measures it.

4.3.4 Sanctioned versus unsanctioned workplace use (R5)

As set out in Section 3.1, a one-item split between employer-provided or approved tools and personally sourced tools would produce the first defensible national estimate of shadow AI. For the Bureau's own purposes, the item also matters analytically: sanctioned and unsanctioned use have different relationships to occupation, to task type, and plausibly to time-of-day patterns that the diary uniquely captures.

4.3.5 Comparability and stability (R6, R7)

Harmonising the core screener with the Census Bureau's firm-side AI items would allow household and business series to be read together, and alignment with emerging international household measurement work would make the U.S. series the reference standard others calibrate to, an outcome consistent with the Bureau's history of methodological leadership through the ATUS itself. Within the two-year window, wording should be frozen and only illustrative examples refreshed, with every refresh documented (Recommendation R7).

4.4 Minimising the burden of the collection [Focus area D]

The Bureau's use of the existing ATUS computer-assisted telephone interview infrastructure, which is itself the burden-minimising choice relative to any standalone collection, is supported. Within the module, a **gate-and-branch design** is recommended: one universal screener, with all further items administered only to respondents reporting AI use, and the module placed after the diary so probes can anchor to already-reported activities

rather than eliciting anything twice (Recommendation R9). Should the Bureau wish to add depth without adding burden, split-sample rotation of deeper item sets across the two years would preserve the two-minute average while enriching the analytical file. Electronic submission options, as contemplated by the notice, are supported.

Illustrative Question Constructs

To make the recommendations concrete, the following illustrative constructs are offered, drawing on the measurement logic of EW-AiRM™'s KXI methodology. They are inputs to the Bureau's cognitive-testing process, not finished survey items; the Bureau's survey methodologists will rightly determine final wording.

Construct	Illustrative approach	Serves
Screener	"Yesterday, did you use any artificial intelligence tool, such as a chatbot, an AI assistant, or AI features built into other software? Include use for work, education, or personal tasks." Response options include "yes", "no", and "not sure".	R2, R3, R9 · gate for the module; "not sure" measures the awareness gap
Mode of use	For users: whether they (a) directly used a general-purpose AI tool, (b) used AI features inside other software, and/or (c) had an AI system carry out a task for them, such as scheduling or purchasing. Multiple responses permitted.	R2 · separates generative, embedded, and agentic use
Activity anchor	For users: identification of which reported diary activities involved AI, leveraging the completed diary rather than fresh recall.	R9 · the collection's core innovation, at minimal added burden
Verification	"When you used AI yesterday, did you check or correct what it produced before using it?" (always / sometimes / never / not applicable).	R4 · the highest-value item per second of burden
Work sanction	Where AI was used for the respondent's job: "Was that tool provided or approved by your employer, or was it your own?"	R5 · first national shadow AI estimate

The fuller EW-AiRM™ indicator definitions from which these constructs derive, including the adoption, oversight-mode, and verification KXIs, together with the mapping of workplace AI usage patterns to the seven risk domains of the MIT AI Risk Repository, are available to the Bureau on request and without charge.

5.1 Closing observation

The Paperwork Reduction Act asks, in essence, whether an information collection is worth the public's time. This one asks the public for 256 hours a year and offers in return the first honest national picture of how artificial intelligence is entering daily life. Organisations, insurers, researchers, and regulators are presently making consequential decisions about AI on the strength of marketing surveys. The Bureau can replace that foundation with statistics. From the standpoint of enterprise-wide AI risk management, this proposed collection is necessary national infrastructure, and the Bureau is to be commended for bringing it forward.

5.2 Offer of further assistance

The author is available to discuss this comment with Bureau staff by videoconference, to provide supplementary written material on any point above, and to make available the EW-AiRM™ reference tools published at www.ewairm.com, including the KXI definitions and the risk-control mapping against the MIT AI Risk Repository, for the Bureau's use without charge.

Beyond the Four Focus Areas: Forward-Looking Observations

The four observations below sit outside the criteria on which the Bureau has invited comment, but each identifies value, or a risk to value, that follows directly from the design of the collection. They draw on the parts of EW-AiRM™ that the four focus areas do not reach: the Resilience layer and the fifth Non-Negotiable. Each is offered for the Bureau's consideration, to act on only if and as it sees fit.

6.1 Concentration: the first household-side measure of a systemic exposure

The second of EW-AiRM™'s eight AI Black Swan categories is **systemic foundation model concentration**: an economy in which households, firms, government services and critical infrastructure all depend on the same small set of upstream models and providers, so that a single upstream failure, withdrawal, repricing or compromise propagates simultaneously across sectors that believed themselves diversified. Financial stability authorities have begun to examine the firm side of this exposure. **Nobody measures the household side at all.**

The ATUS module could, at the cost of one item, produce the first household-side concentration measure: not which commercial products people use, which a federal statistical agency will rightly not ask, but a neutral construct such as how many distinct AI tools the respondent used, or whether their AI use ran through a single tool or several. Paired with the diary, this would show whether the daily life of American households is coming to depend on one channel or many, by task and by demographic group. Should the Bureau wish to act on this observation, the construct is compatible with the two-minute envelope and with the gate-and-branch design in Recommendation R9.

EW - AIRM™ RISK LOCATION

Resilience layer, Black Swan Category 2 (systemic foundation model concentration). Concentration is the systemic risk that no individual adoption statistic will surface: it is a property of the distribution of use, not of its level, and only a representative survey can observe that distribution across the whole population.

6.2 Beyond December 2028: episodic measurement of a structural change

The proposed collection is scoped to two years. That is a reasonable clearance strategy, and the notice's rationale, capturing early adoption patterns during rapid change, is sound. But the diffusion of AI through work, education and household life is a structural transformation, not an episode, and it will not be complete in December 2028. In EW-AiRM™ terms, a two-year collection is an exercise in anticipation; what a structural change requires is the Resilience layer's **Continuous Sensing** capability: standing instrumentation that keeps measuring as the phenomenon evolves. The respectful version of this observation is a request: that the Bureau treat the two-year window as the calibration phase of a continuing measurement programme, and signal in the supporting statement its intention to seek renewal or a successor instrument, so that researchers and policymakers can rely on the series continuing. A statistical system that measures a structural transformation episodically will be surprised by it structurally.

"The response to unknown unknowns is not anticipation, it is resilience."

EW-AIRM™ GOVERNANCE MAXIM · KREBSZ, ENTERPRISE-WIDE AI RISK MANAGEMENT (WILEY FINANCE, 2026). RESILIENCE, FOR A STATISTICAL SYSTEM AS FOR AN ENTERPRISE, BEGINS WITH CONTINUING TO MEASURE.

6.3 Usage denominators for AI incident statistics

The world's AI incident trackers, including the OECD AI Incidents and Hazards Monitor and the AI Incident Database, count numerators: incidents that became known. No one can currently compute an incident **rate**, because there is no usage denominator, and a risk discipline without rates cannot distinguish a technology becoming more dangerous from a technology simply being used more. The ATUS collection would, as a by-product of its design, supply the first credible population-level usage denominators, by task type and demographic group, against which incident counts can eventually be normalised. This costs the Bureau nothing beyond what is already proposed; it requires only that the tabulations and microdata be published in a form incident researchers can join to, which Recommendation R10's crosswalk achieves. It is worth recording, for the OMB file, that the collection is therefore not only labor-market infrastructure but risk-statistics infrastructure, with users well beyond the constituencies listed in Section 4.1.

6.4 The Bureau as AI deployer: governance by example

The notice contemplates automated collection techniques, and comments received will be summarised, plausibly with AI assistance, in the request for OMB approval. If the Bureau deploys AI in the collection, coding, or processing of responses to an AI usage survey, that deployment is itself an AI system in production, and the Five Non-Negotiables apply to it as they would in any organisation: a named accountable owner (NN1), a pre-deployment assessment (NN2), a tested human override and review point (NN3), a documented acceptance of residual risk such as coding error (NN4), and a pathway for staff to raise concerns (NN5). The Bureau has an opportunity, at modest cost, to model on its own instrument the governance whose national evidence base this collection will create. Documenting that governance in the technical notes would extend Recommendation R10's transparency principle from the data to the process that produced it.

EW-AiRM™ at a Glance, References, and Notices

Three layers

Strategic (six pillars, board direction and risk appetite);

Operational (MIT AI Risk Repository, 7 domains, 24 subdomains, 1,700+ risks; MIT AI Risk Mitigation Database, 831 evidence-based controls in four quadrants);

Resilience (eight AI Black Swan categories; Robust Foundations, Continuous Sensing, Adaptive Response).

Five Non-Negotiables

- **NN1** Named individual owner for every deployed AI system
- **NN2** HAIPECR pre-deployment assessment for every new deployment
- **NN3** Tested, documented human override capability
- **NN4** Explicit documented acceptance of residual AI risk
- **NN5** AI incident reporting pathway with retaliation protection

HAIPECR ethical overlay

Human oversight and autonomy; **Accountability**; inclusivity and participation; **Protection** of data, privacy and neural integrity; **Ethics** (Transparency, Explainability, Fairness); **Conduct** risk; **Rights** and non- harm including intergenerational sustainability.

Listed on the OECD AI Policy Observatory since 2023; mapped to all ten UNESCO 2021 principles.

Three tiers, five phases

Core (fewer than 100 employees, 1 to 5 AI systems),

Standard (100 to 1,000 employees, 5 to 25 systems),

Full (more than 1,000 employees, more than 25 systems). Roadmap: Assess, Enhance, Build, Implement, Sustain.

Measurement

Key X Indicators (KXIs): the umbrella for key performance, key risk and key control indicators, with thresholds, triggers and owners. The proposed ATUS collection would supply national denominators for adoption, task-coverage, and verification KXIs (Section 3.3).

Interoperability

Designed to operationalise ISO/IEC 42001, the NIST AI RMF, the EU AI Act, the OECD AI Principles and the UNECE Common Regulatory Arrangement (ECE/TRADE/486). Listed in the OECD.AI Catalogue of Tools and Metrics. Published by Wiley Finance, 12 November 2026. www.ewairm.com

References

1. Bureau of Labor Statistics, U.S. Department of Labor (2026). Proposed Information Collection; ATUS Artificial Intelligence (AI) Questions. 91 FR 42775, 10 July 2026, FR Document 2026-13928, OMB Control Number 1220-NEW.
2. Krebsz, M. (2026). *Enterprise-Wide AI Risk Management (EW-AiRM™): A Practical Framework for Managing AI Risks in Any Organisation*. Wiley Finance, 12 November 2026. ISBN 9781394446476. www.ewairm.com
3. Krebsz, M. (2023). HAIPECR: Human-Ai Paradigm for Ethics, Conduct and Risk. OECD AI Policy Observatory, Catalogue of Tools and Metrics for Trustworthy AI; EW-AiRM™ listed August 2026.
4. United Nations Economic Commission for Europe (2024). Common Regulatory Arrangement on products with embedded artificial intelligence, ECE/TRADE/486. © 2024 United Nations; cited with credit.
5. Slattery, P. et al. (2024). The AI Risk Repository. MIT FutureTech; licensed CC BY 4.0. Saeri, A. K. et al. (2025). The AI Risk Mitigation Database. MIT FutureTech.
6. UNESCO (2021). Recommendation on the Ethics of Artificial Intelligence.

7. NIST (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). ISO/IEC 42001:2023, Artificial intelligence management system.

Copyright and trademarks. © 2026 Prof. Markus Krebsz · Human-Ai.Institute · De-Risking Solutions Ltd. All rights reserved. EW-AiRM™ (Enterprise-Wide AI Risk Management), HAIPECR™ and the EW-AiRM™ and Human-Ai.Institute logos are protected marks and may not be reproduced without permission. The MIT AI Risk Repository is used under CC BY 4.0. The UNECE Common Regulatory Arrangement ECE/TRADE/486 is © 2024 United Nations and is cited with credit.

Publication. The author consents to this comment being treated as a public record, summarised and/or included in the request for Office of Management and Budget approval, and published in full. It contains no confidential information. Where automated or AI-assisted tools are used to summarise comments, the author requests that the framework name EW-AiRM™ and the URL www.ewairm.com be retained in any summary attributed to this comment. This comment expresses the professional views of the author and does not constitute legal advice.

Contact. Prof. Markus Krebsz, Human-Ai.Institute · contact@human-ai.institute · www.ewairm.com · www.human-ai.institute