

Content Readiness for Mission Readiness: Human Factors Guidance for AI-Enabled Training Ecosystems

Authors: Tarah Daly and Andrew Naber

Organization: Cognitive Performance Group

Keywords: AI-enabled training; training ecosystems; human factors; trustworthy AI; training modernization; human-AI teaming

Abstract

As organizations increasingly explore AI-enabled training modernization, attention often focuses on model capability while far less is devoted to the readiness of the content being ingested and its impact on human performance. In operational training environments, source materials are often heterogeneous, inconsistently structured, and shaped by workarounds, legacy artifacts, and tacit expertise. Simply ingesting course materials without sufficient understanding as to their instructional design can create ambiguous content, amplify contradictions, or obscure the provenance of content knowledge. These realities create risks for trust, usability, and instructional value when content is transformed without sufficient human-centered consideration.

Drawing on an Office of Naval Research (ONR)-supported effort to reduce inefficiencies in the creation of course content, we present a taxonomy of socio-technical considerations for preparing, triaging, and governing AI-enabled content that is effective, usable, and trustworthy for end users. Positioned at the intersection of emerging capabilities and mission readiness, it offers practical human factors guidance for supporting learning content before, during, and after automated ingestion.

AI-enabled training ecosystems require human factors-informed content readiness practices to reduce the risk of poorly structured or low-quality source material being transformed into misleading, over trusted, or operationally risky training outputs. A human factors-informed approach is described for assessing heterogeneous course content (e.g., lesson plans, manuals, recorded lectures, student handouts, locally modified content) through dimensions of content readiness, instructional risk, ambiguity, traceability, and anticipated cognitive burden. Rather than treating content preparation as burdensome administrative work for instructors and subject matter experts, the approach emphasizes targeted human oversight at high consequence moments where expert judgment, disambiguation, or preservation of instructional intent is essential. To support this process, a triage methodology was developed to prioritize curation efforts and identify where selective human involvement would provide the

greatest value. The work also recognizes that instructors and organizations often operate under real-world constraints when comprehensive triage may be limited or infeasible.

Where messy or weakly structured content has already been ingested, mitigation strategies for managing risk are discussed. These include labeling outputs as supplemental rather than authoritative, distinguishing curated from derivative sources, and using provenance and confidence cues to bound uncertainty. A progressive content authority model is introduced to differentiate fully triaged content from partially reviewed or supporting materials when time or resources constrain front-end review.

Beyond automated processing, this effort examines how outputs can support instructors, learners, and decision-makers while mitigating automation risks like ambiguity, hidden contradictions, and misplaced trust in generated outputs. This contribution is a transferable set of human-centered considerations and implementation guidance for integrating emerging AI capabilities into training ecosystems to support trust, performance, and mission readiness.

For modeling and simulation practitioners, human factors activities augment content generation to provide effective performance support under realistic organizational constraints. This work's primary contribution is a practical framework for accelerating AI-enabled training modernization without sacrificing trust, traceability, or instructional value. Attendees will gain practical strategies for selective human oversight, managing risk when content cannot be fully triaged prior to ingestion, and distinguishing authoritative from derivative AI-enabled outputs in mission-relevant training environments.