

UXLS Community

Workstream: UX for AI and AI for UX
Quick start learning resources

Autumn 2025



UX for AI – An ongoing conversation

This is a resource pack assembled by the UXLS AI workstream to enable UX practitioners working in the Life Sciences to get some basic understanding across three different areas:

- **When is an AI solution appropriate**
- **Agentic AI**
- **Ethical AI and trust**

This resource pack was assembled in Q1 and Q2 of 2025.

What this isn't

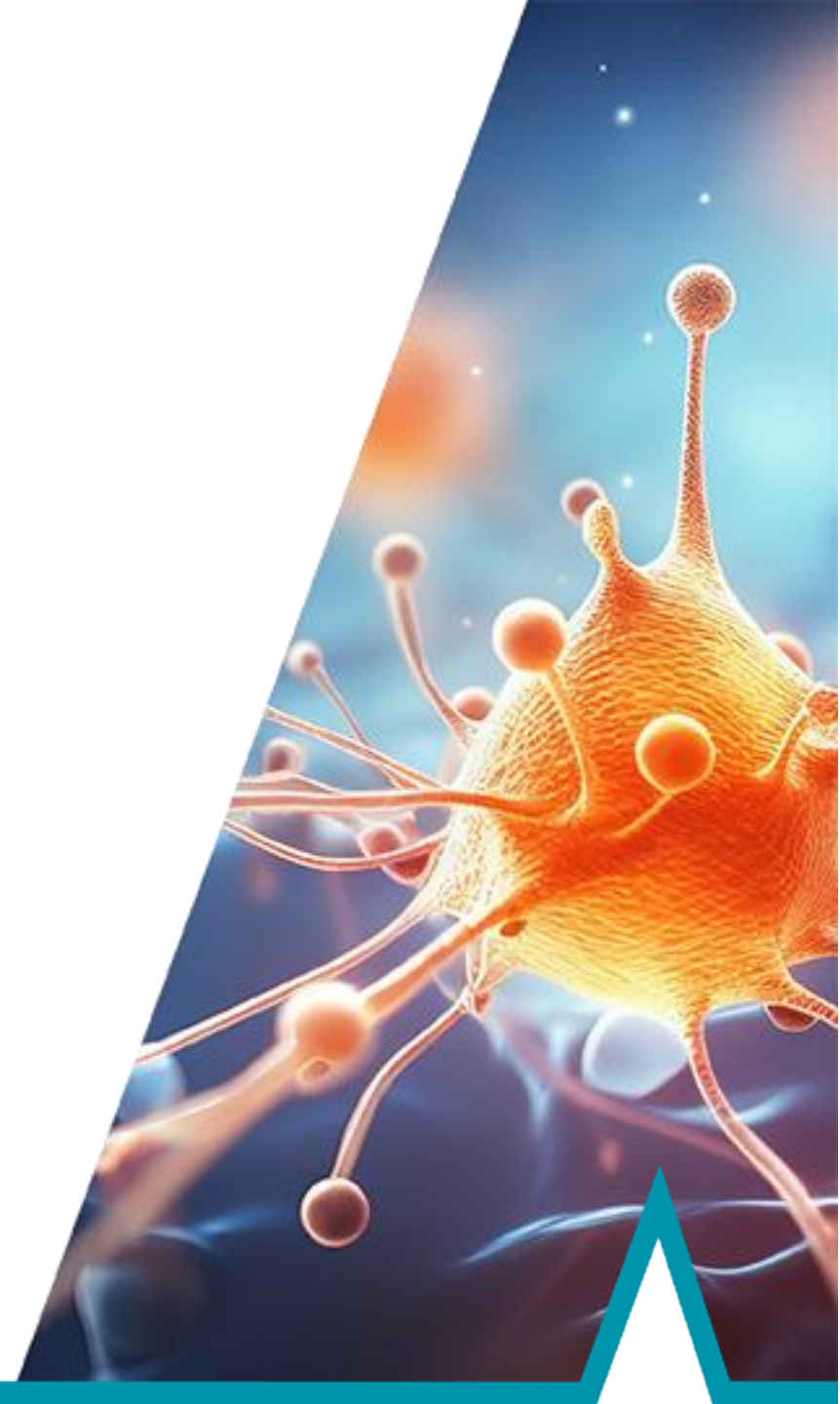
A complete training course on AI - but rather a starting point to some resources we have found useful.



When is an AI solution appropriate?

Resources that cover:

1. When to use AI?
2. What flavour of AI to use?

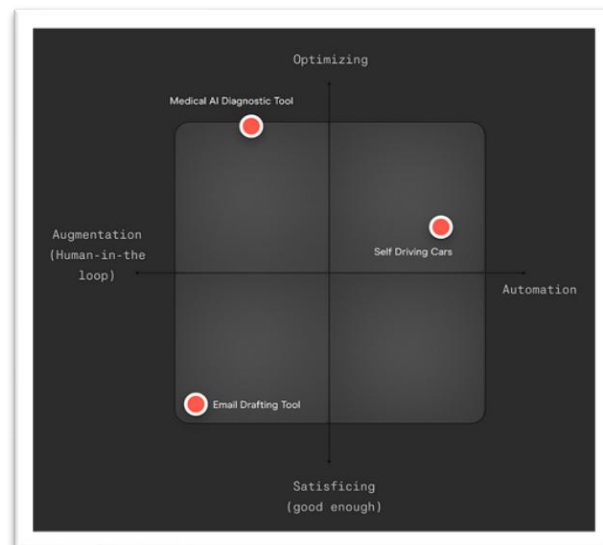


Do you need AI to solve the problem?

Useful resources to help you and your team identify which AI opportunities are the most appropriate ones for your problem space

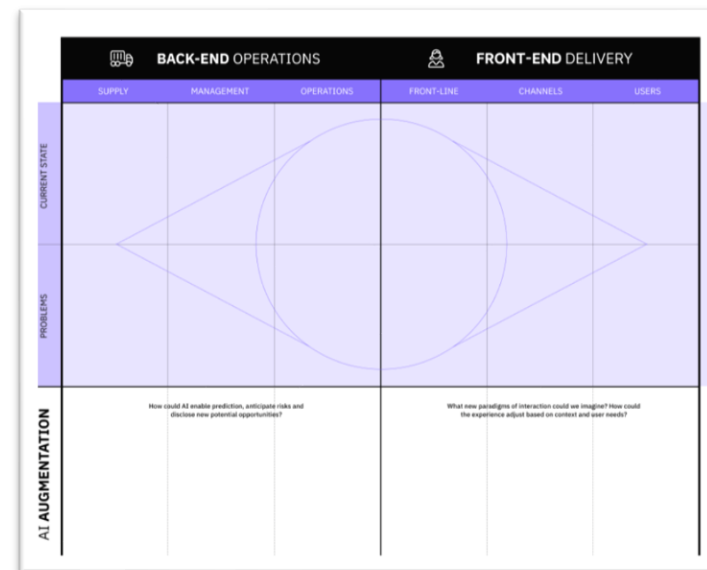
Designing AI with purpose: the AI intention matrix by Richard Yang

A framework to help teams manage the (excessive) enthusiasm to push for advanced AI features when they don't add much value to the users



AI opportunity landscape by Oblo

A framework to help teams identify which AI opportunities can be integrated based on the understanding of user needs and inspired by the technological capabilities.



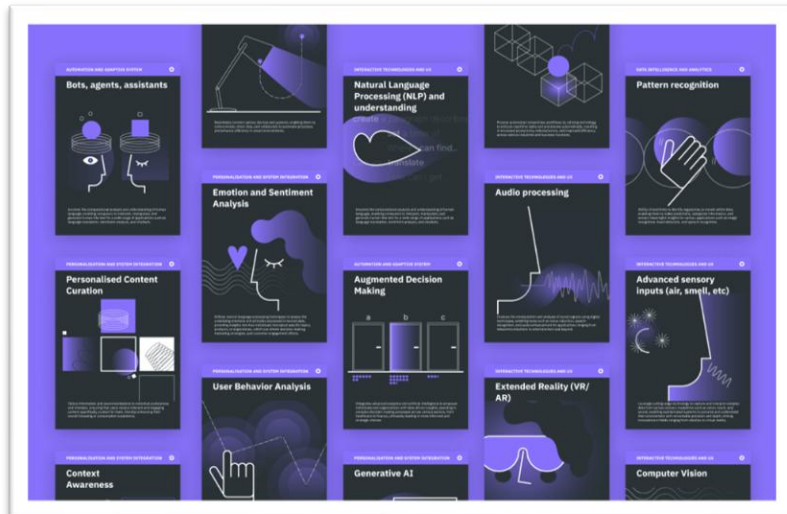
What AI capabilities do you need?

Useful resources to help you and your team identify which specific AI capabilities are most appropriate for your problem space

AI functionality cards

by Service Design tools

A set of cards of current AI capabilities to stimulate the ideation process and inform the generation of opportunities for AI in a given problem space



AI cards

by 33A

AI capability cards to help you and your team to decide where to apply AI in a given setting



AI Design Patterns

AI Design Pattern Framework for AI/UX

by Dave Brown

Framework to guide designers to create effective and intuitive AI experiences design principles, considerations and examples of how to address them through five different phases:

- **AI onboarding:** Expectation framing and intent scaffolding
- **User input:** Support users in creating the right inputs for AI
- **AI output:** Design good processing cues, confidence visualisation and explainability and how to support users with multiple outputs from AI
- **Undo and version control:** Design patterns to allow users to correct, reverse, or refine interactions
- **System learning:** Patterns to design how the AI system adapts, learns and improves over time

Expectation Framing

What it solves:

Users often assume software is deterministic. In AI systems, this can lead to misplaced trust or frustration.

Framing expectations early — before the user even interacts — helps establish a healthy model of the system's capabilities and limitations.

Why it matters:

Users bring mental models from traditional software into AI-driven interfaces.

Without upfront messaging that frames the system as collaborative, probabilistic, and imperfect, users may expect precision and get frustrated by unexpected or "wrong" behavior.

Design Considerations:

- How do we signal that this is a collaborative system, not an oracle?
- Can we use friendly, human-centered language to build trust without overpromising?
- Are we clear about the kinds of things the AI is good and bad at?

Examples:

Product	Example Phrase	Context Location
GitHub Copilot	"I can help, but I make mistakes."	Welcome tooltip/copy
LinkedIn	"People you may know"	Section headers
ChatGPT	"May occasionally produce inaccurate results"	System message or UI bar

Expressive Input

What it solves:

AI outputs require the right input. Yet, users may struggle to express nuance, tone, or creative intent through traditional natural-language prompts alone.

Expressive input patterns provide lightweight, visual, or emotionally rich ways to shape AI behavior — enhancing intent clarity or creative exploration without requiring detailed language.

Why it matters:

These patterns make AI feel more approachable, fun, and intuitive. They invite experimentation, support faster iteration, and reduce prompt anxiety — especially for users who don't know what to type.

Expressive inputs can also unlock more personalized, emotionally resonant experiences without adding cognitive load. These patterns are especially helpful in creative, multimodal, or mobile-first environments.

Design Considerations:

- Are expressive inputs intuitive, or do they need onboarding?
- Do users feel in control or surprised by how the AI interprets them?
- Can these inputs be combined with traditional prompts?
- Are expressiveness tools accessible across modalities (text, voice, visuals)?
- Is it clear how the system is interpreting visual/emotional cues?

Examples:

Product	Input type	Context
Midjourney	Emoji-as-prompt	Visual generation input
Luma AI	Madlib-style prompt builder	Encourages structured creativity
Figma	2x2 grid	Modifying copy across a range of options
Canva Magic Write	Icon-based prompt builder	Guided content creation

Processing

What it solves:

AI responses often involve invisible, multi-step operations — from calling APIs or MCPs, to generating drafts or running retrievals. Without clear signals that work is happening, users may feel confused, impatient, or assume the system is broken.

Processing cues reassure users that the system is active and build trust by revealing how work is being done, not just that something is coming.

Why it matters:

Good processing design shapes user expectations around time, effort, and trust. It manages impatience, prevents premature abandonment, and builds transparency — especially in agentic or tool-using systems where multiple steps are hidden. It also allows users to calibrate mental models: "Is this AI just thinking, or is it stuck?"

Design Considerations:

- Is it clear that the system is working — not stalled — through animations or visible cues?
- Is the wait time contextualized appropriately, with progress indicators or previews for longer tasks?
- Is transparency about processing balanced with cognitive load, depending on task complexity?

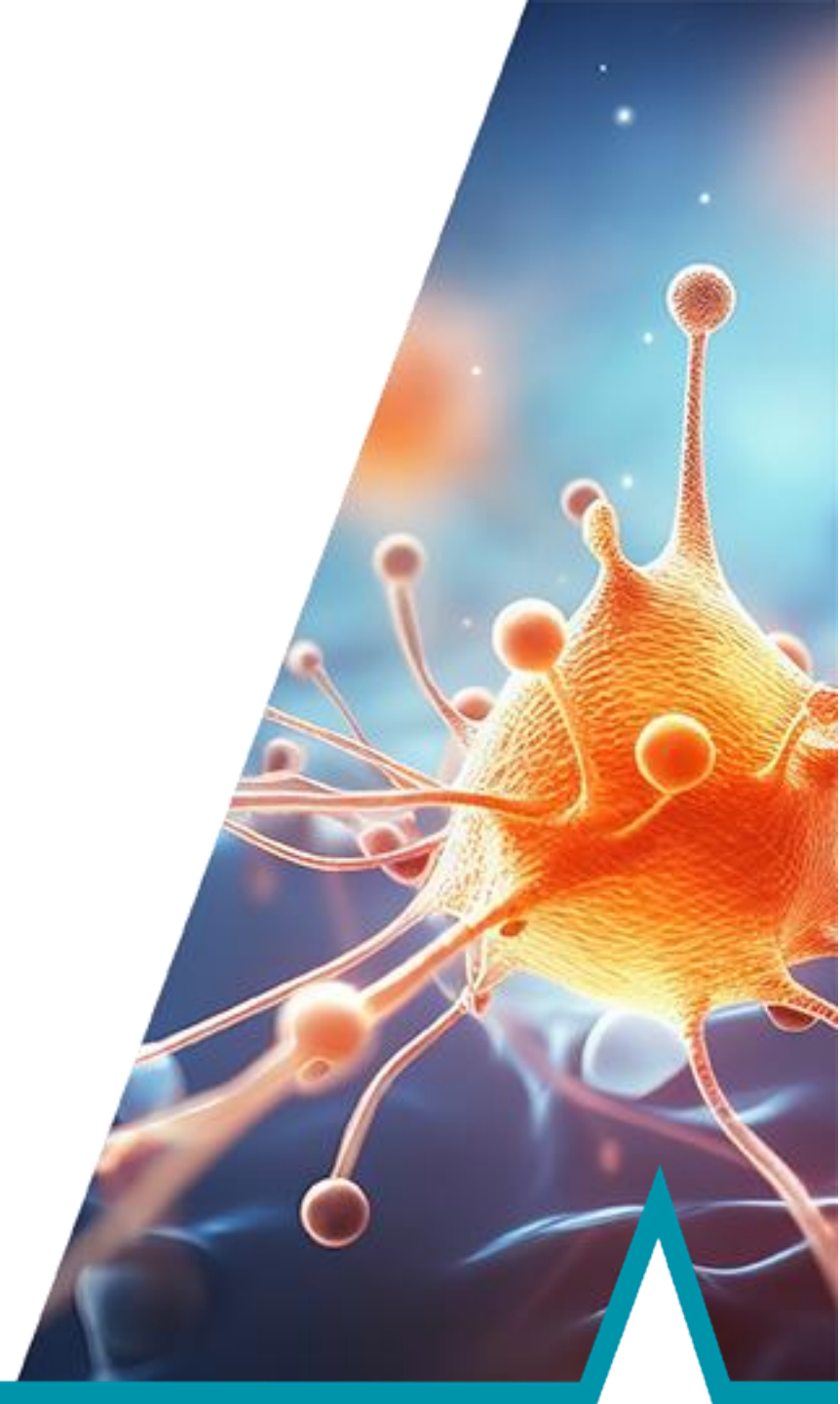
Examples:

Product	Processing mechanism	Context	Type
ChatGPT	Typing dots animation (..) while streaming response	Chat generation — shows active thinking	Streaming
Claude (Artifacts)	Typing animation while drafting, live updating doc preview	Full document creation with inline artifacts	Streaming + partial steps
Illicit	Show retrieval, search, and reasoning stages visibly before synthesis	Research agent surfacing work-in-progress	Processing Steps
Attio AI	Status updates like "Analyzing meeting notes..." → "Drafting summary..."	CRM workflow with multi-stage processing	Processing Steps

Agentic AI

Topics covered:

1. What is Agentic AI?
2. Agentic AI resources



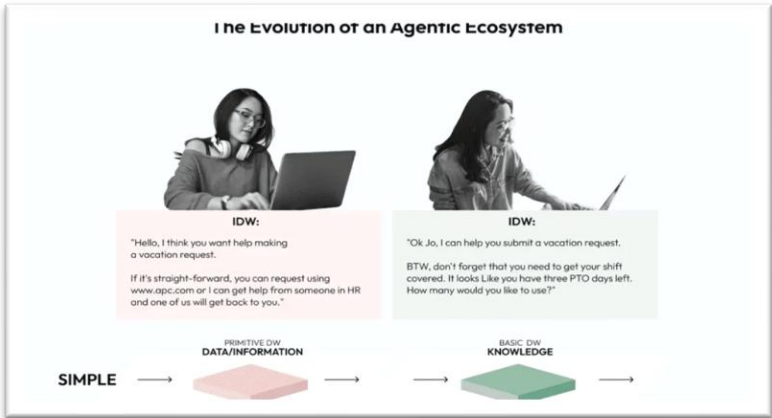
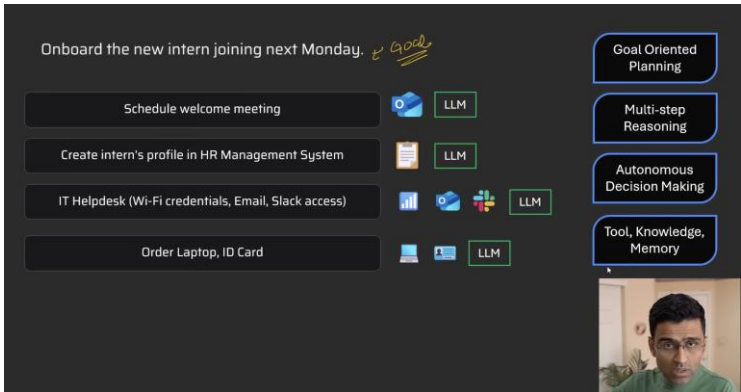
What is Agentic AI ?

We recommend the following videos that give a basic introduction to Agentic AI and provide a comparison of Agentic AI versus Generative AI with some excellent examples:

What is Agentic AI and How Does it Work?
by Codebasics

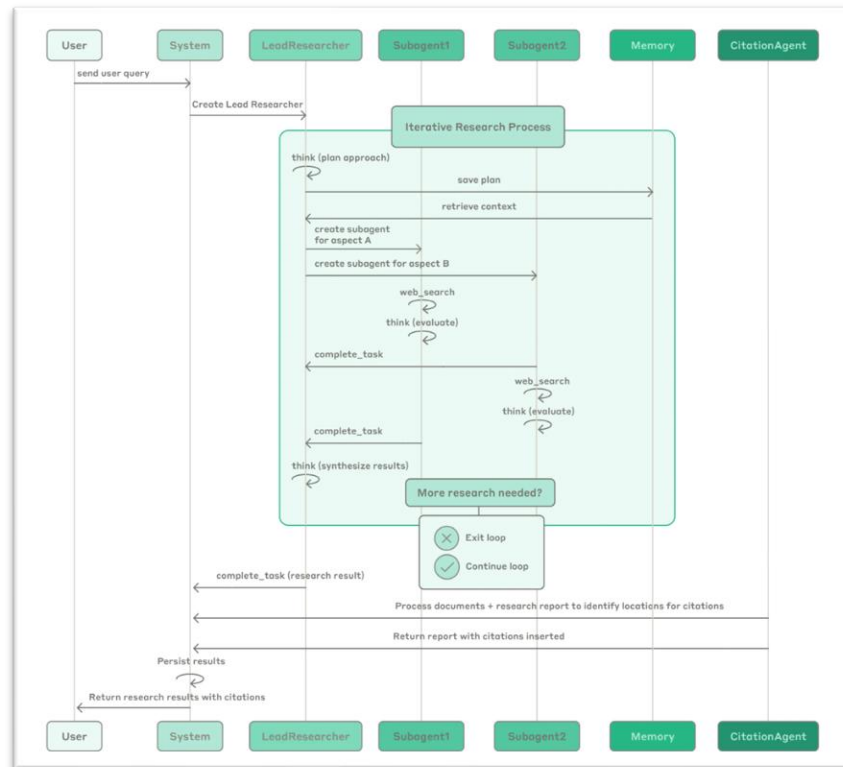
Generative AI vs Agentic AI: Shaping the future of AI collaboration
by IBM Technology

Agentic AI: Fostering Autonomous Decision Making in the Enterprise
by UX Magazine

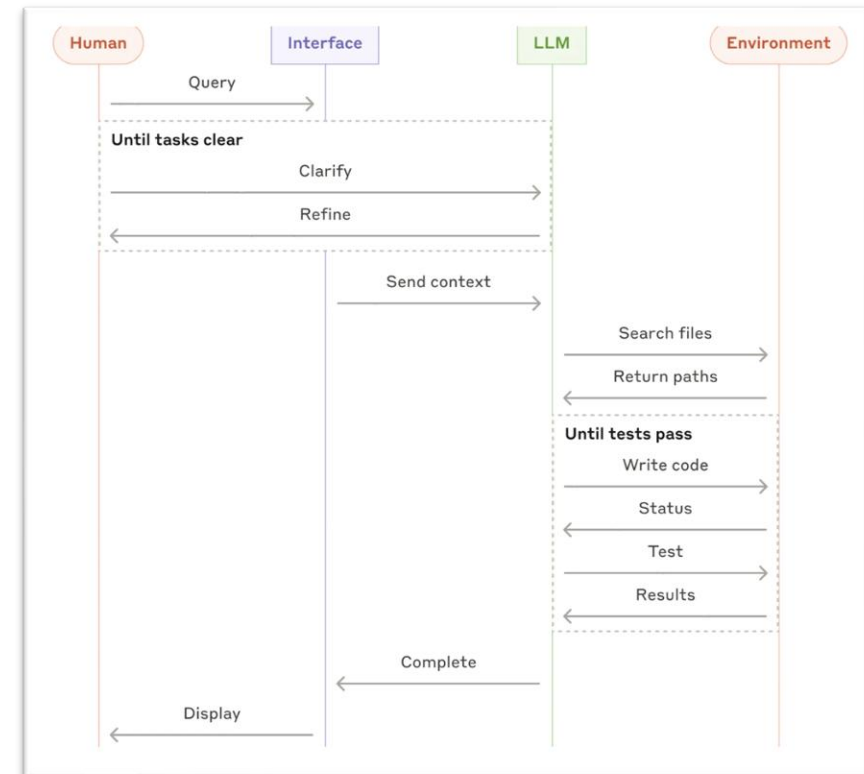


Agentic AI additional resources

Case study on how to create a multi-agent research system by Anthropic



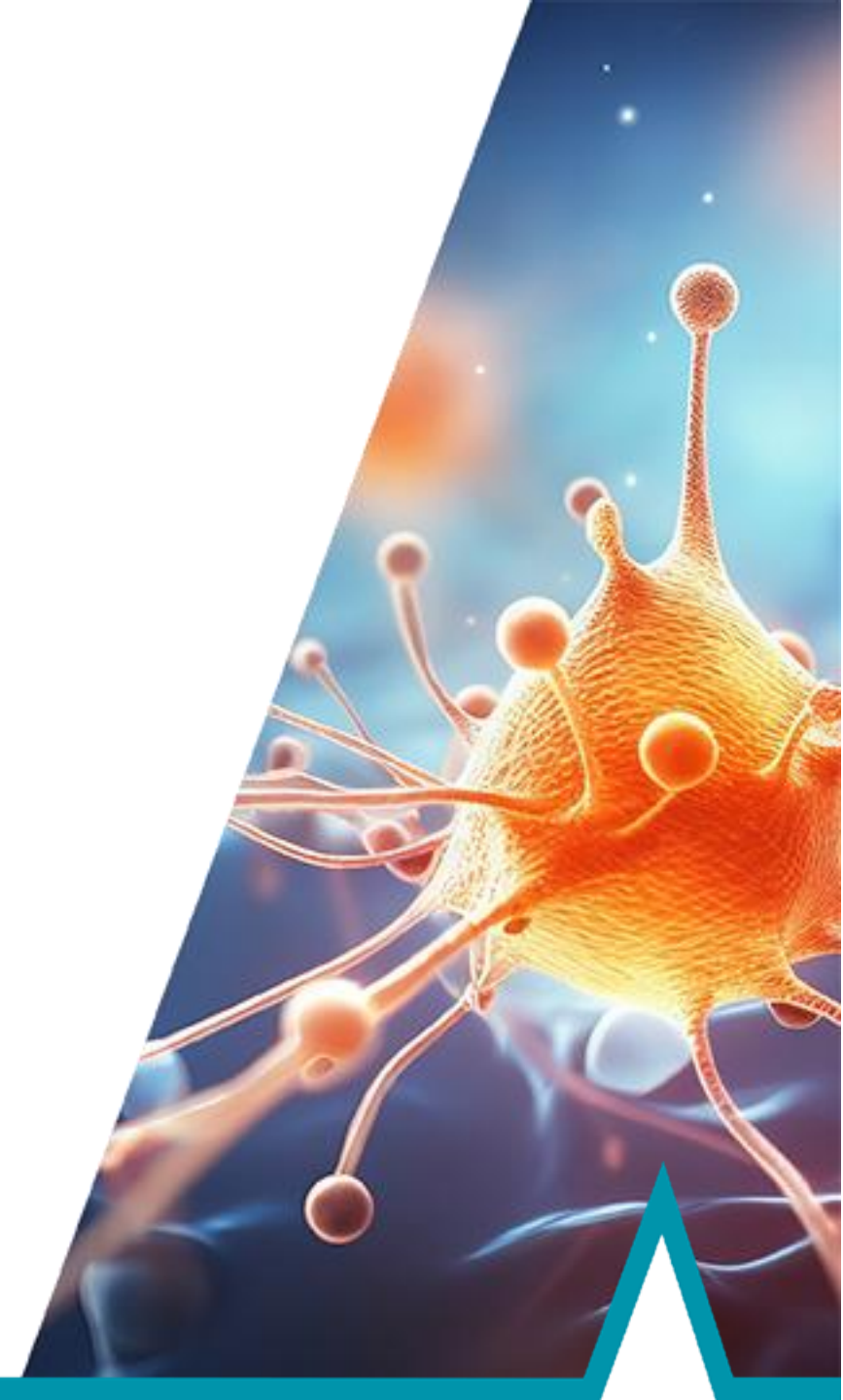
Guidelines on building effective agents by Anthropic



Legal and ethical responsibilities

Topics covered:

1. Legal responsibilities and the EU AI act
2. Ethical responsibilities
3. GxP regulated environments

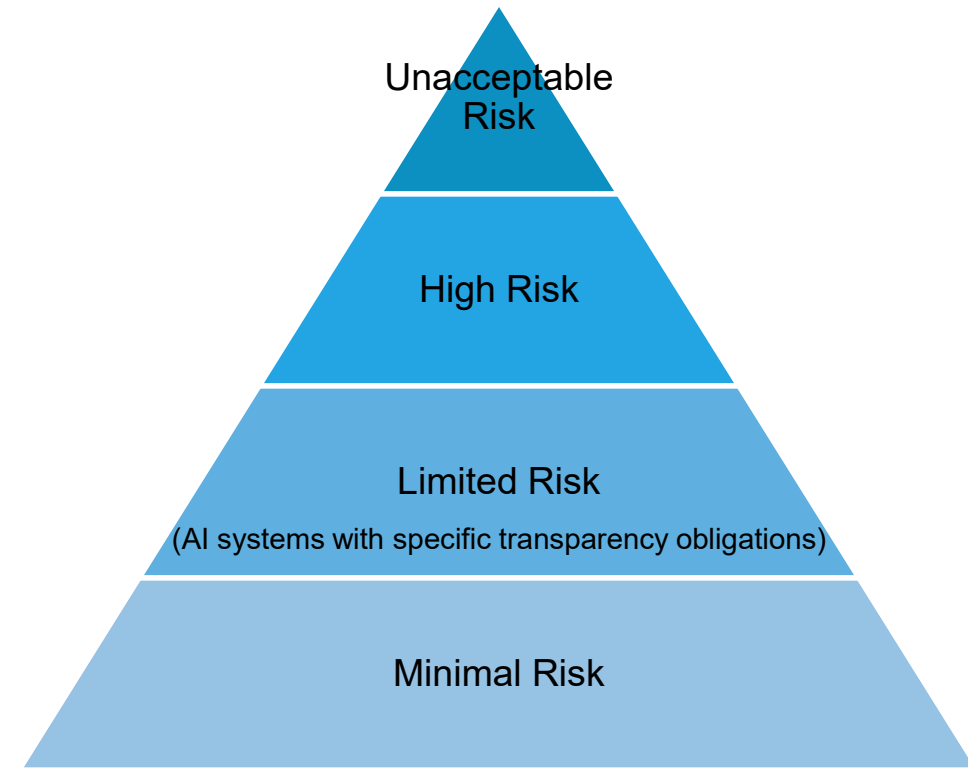


Legal responsibilities

EU AI Act (2024)

The EU AI Act introduces a uniform framework based on a forward-looking definition of AI and a risk-based approach:

- **Minimal risk:** most AI systems such as spam filters and AI-enabled video games face no obligation under the AI Act, but companies can voluntarily adopt additional codes of conduct.
- **Specific transparency risk:** systems like chatbots must clearly inform users that they are interacting with a machine, while certain AI-generated content must be labelled as such.
- **High risk:** high-risk AI systems such as AI-based medical software or AI systems used for recruitment must comply with strict requirements
- **Unacceptable risk:** for example, AI systems that allow “social scoring” by governments or companies are considered a clear threat to people's fundamental rights and are therefore banned.



The design should accommodate for the level of risk identified and minimize this risk where possible.

Legal responsibilities

Further information on the EU AI Act (2024)

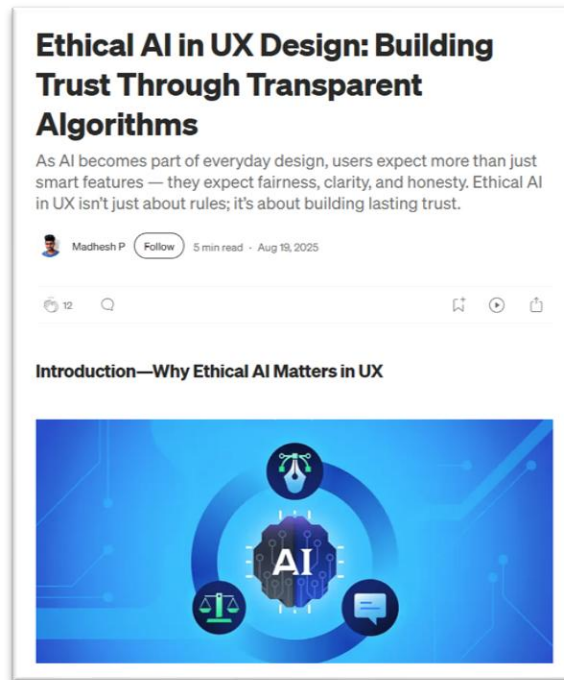


[EU AI act video](#)

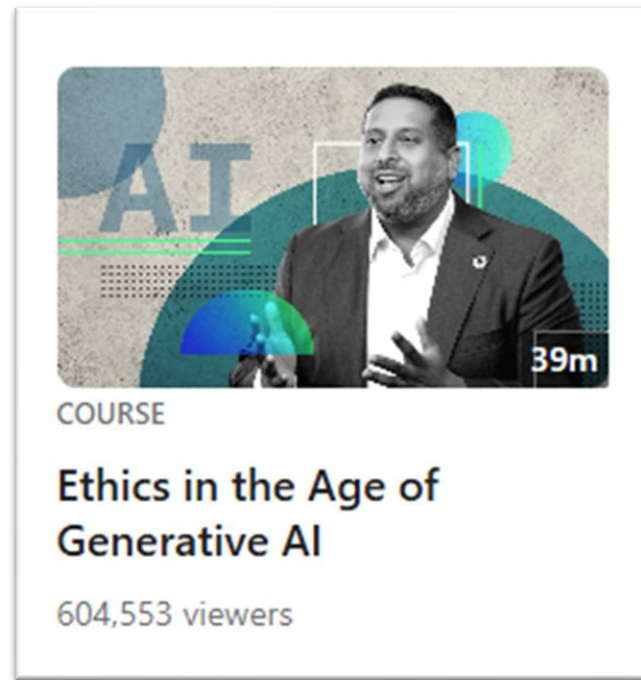
A screenshot of the European Commission website page for the AI Act. The page has a blue header with the European Commission logo and the text "Shaping Europe's digital future". Below the header is a navigation bar with links: Home, Policies, Activities, News, Library, Funding, Calendar, Consultations, and AI Office. The main content area is titled "AI Act" and contains a "PAGE CONTENTS" sidebar on the left with links to various sections: "Why do we need rules on AI?", "A risk-based approach", "How does it all work in practice for providers of high-risk AI systems?", "A solution for the trustworthy use of large AI models", "Supporting compliance", and "Governance and". The main text area on the right provides an overview of the AI Act, stating it is the first-ever legal framework on AI, addresses risks, and positions Europe to play a leading role globally. It also mentions the AI Act (Regulation (EU) 2024/1689) and the AI Pact. A "Share" button is located to the right of the main text. A "Quick links" sidebar on the right contains links to the AI Act Single Information Platform, The AI pact, Text of the AI Act in all EU official languages, Impact assessment of the regulation, Study supporting the impact assessment, FAQs: New rules for Artificial Intelligence, and European AI Office.

Read more: <https://digital-strategy.ec.europa.eu/en/policies/regulatory-framework-ai>

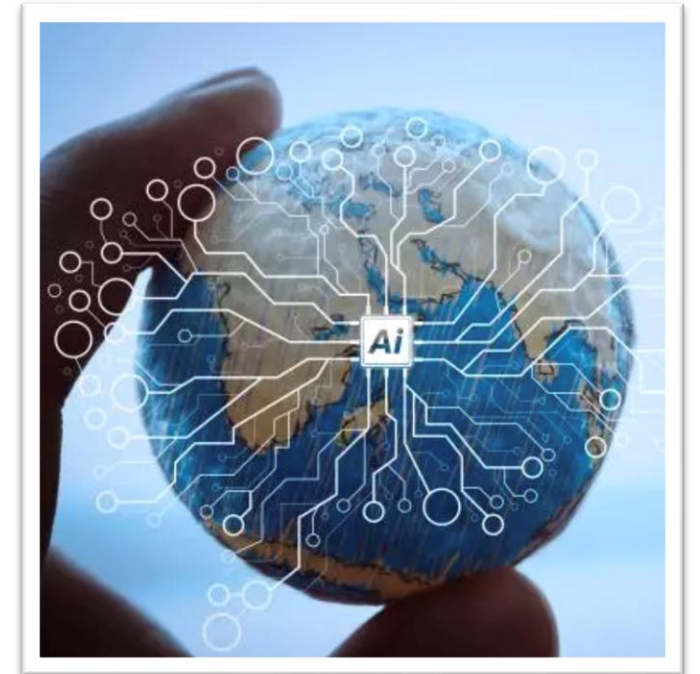
Ethical responsibilities



Ethical AI in UX design
by MadPesh P



Ethics in the Age of Gen
AI
by Vilas Dhar



UNESCO: Ethics of
Artificial Intelligence



GXP regulated environments – Annex 22

- Annex 22 is a newly introduced section in EudraLex Volume 4 (EU GMP Guide), specifically addressing the use of Artificial Intelligence (AI) and Machine Learning (ML) in pharmaceutical and biologics manufacturing and quality systems
- Annex 22 provides guidelines of AI usage in a GXP environment.
- It formally acknowledges AI/ML in GMP contexts—particularly for static, deterministic models with direct impact on patient safety, product quality, or data integrity

2. Principles

- 2.1. *Personnel.* In order to adequately understand the intended use and the associated risks of the application of an AI model in a GMP environment, there should be close cooperation between all relevant parties during algorithm selection, and model training, validation, testing and operation. This includes but may not be limited to process subject matter experts (SMEs), QA, data scientists, IT, and consultants. All personnel should have adequate qualifications, defined responsibilities and appropriate level of access.
- 2.2. *Documentation.* Documentation for activities described in this section should be available and reviewed by the regulated user irrespective of whether a model is trained, validated and tested in-house or whether it is provided by a supplier or service provider.
- 2.3. *Quality Risk Management* Activities described in this document should be implemented based on the risk to patient safety, product quality and data integrity.

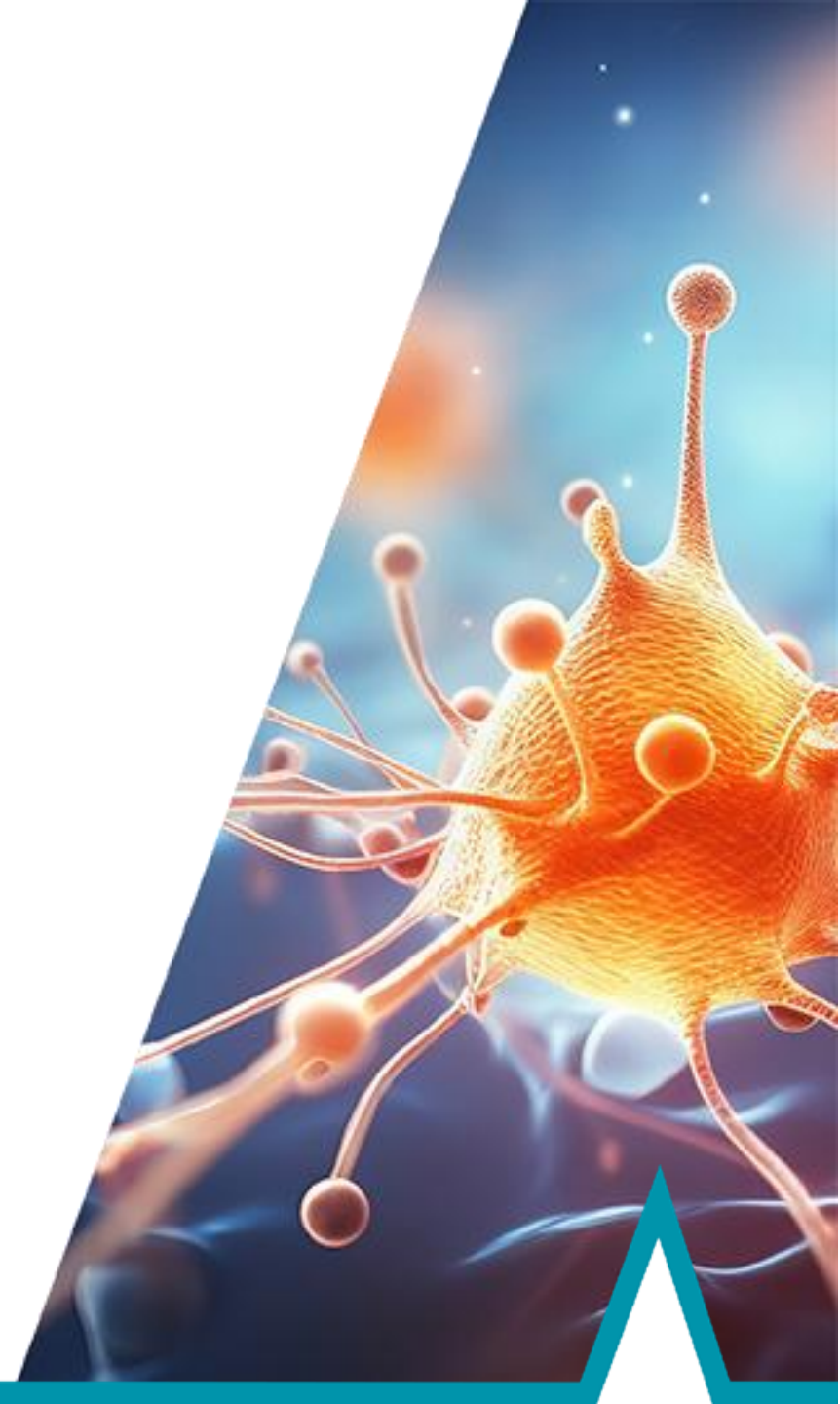
Read more on the [Annex 22](#)



Evaluating AI

Topics covered:

1. Bot usability scale



Measuring Chatbots usability

Currently our workstream have not found many useful resources to measure usability of AI.

We have found mentions of chat bot usability through the Bot Usability Scale below.

Bot Usability Scale (BUS)

- Borsci, S., Malizia, A., Schmettow, M. *et al.* The Chatbot Usability Scale: the Design and Pilot of a Usability Scale for Interaction with AI-Based Conversational Agents. *Pers Ubiquit Comput* **26**, 95–119 (2022).
<https://doi.org/10.1007/s00779-021-01582-9>

Although not specifically measuring trust the authors have identified the following to be true:

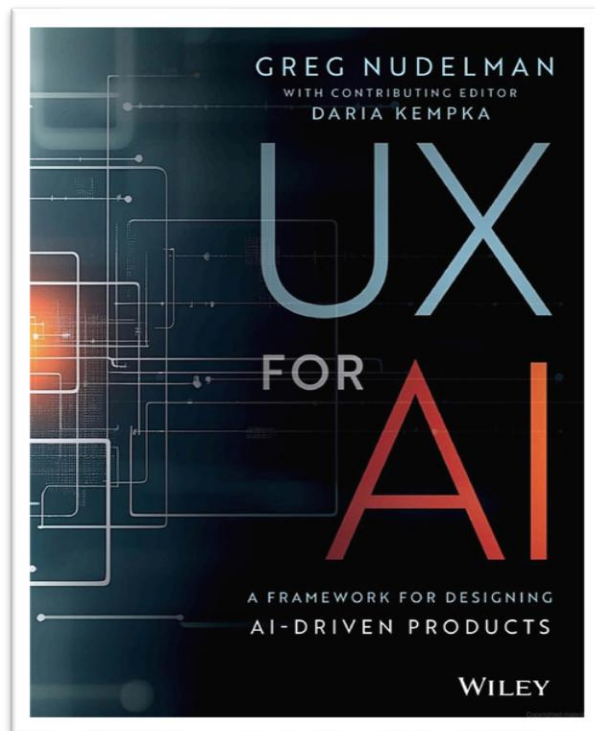
- *"It is easier to assess 'trust' in a CRM chatbot interaction by assessing the bot's capacity to provide information and helping to attain a goal (i.e. the credibility of information) instead of by assessing trustworthiness as a general and unspecified sense of trust. Assessing trustworthiness could require a different set of items more in line with trust and technology acceptance theory"*



Useful resources recommended by our workstream

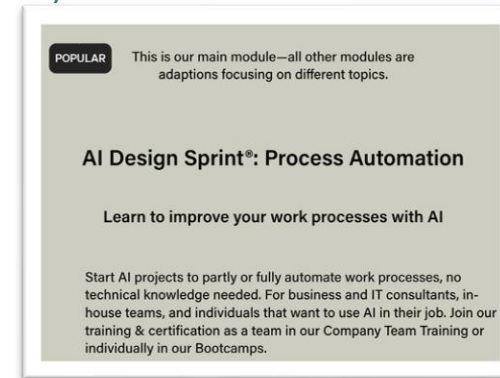
Books:

[UX for AI: A Framework for Designing AI-Driven Products](#) by Greg Nudelman

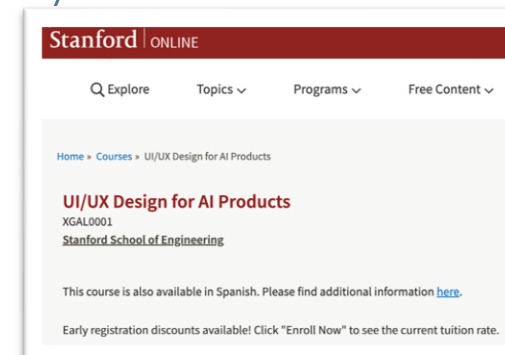


Courses:

[AI Design Sprint courses](#) by 33a.ai



[UI/UX Design for AI products](#) by Stanford



Resource created in 2025 by

AI for UX and UX for AI workstream

Pistoia Alliance UXLS Community

Workstream leads

Voula Gkatzidou (GSK)

Paula de Matos (AVEVA)

Workstream members

Peter Hummel (Novartis)

Nelson Taruc (Lextech)

Roger Attrill (IQVIA)

Sudha Yerramilli (Merck)

Fernando Barneze Gonzalez
(Novartis)

Julie Morrison (Instem)

Breac Baker (Instem)

Brian Mila (Lextech)

Kasia Konczak (GSK)

Aleksandra Zajdenc-Jurczyk
(Roche)

John Wise (Pistoia Alliance)

Anne Stevens (Elsevier)

Jing Zhang (AstraZeneca)

Adriana Rys (Novartis)

Ahmet Bektes (Elsevier)



Join the journey



[Pistoiaalliance.org/community/uxls](https://pistoiaalliance.org/community/uxls)

Community Facilitator: Farah Egby
farah.egby@pistoiaalliance.org

UXLS

User Experience for Life Sciences

COMMUNITY

PISTOIA
ALLIANCE



Our members
collaborate as equals
on open projects that
generate significant
value for the worldwide
life science community

PISTOIA ALLIANCE

UXLS Community Funders



Thanks to our funders who are
making this possible, without
their help this community
would not exist or collaborate.
Thanks for all the in-kind
contributions for their time
and expertise