

AICON

A field guide to the ideas

Understand the concepts. Follow the connections. Find the material worth your time.

Prepared by Acuity AI · 12 September 2026

An independent learning resource based on the announced programme and selected public material. Not an official conference publication.

Use this guide to prepare for sessions, explore an unfamiliar concept, or bring a better question back to your organisation. Each chapter includes selected reading, programme connections and related ideas.

Reading recommendations and relationships between concepts are editorial synthesis. Conference connections refer to announced roles and sessions, not confirmed attendance or endorsement. Video review limits are stated beside each resource.

acuityai.co/AICON

Find your starting point

Knowledge readiness & authority

For one important business question, which source wins when two documents disagree, and who decides?

Retrieval-augmented generation

Can the system retrieve the passage that would let a knowledgeable colleague answer this question?

Knowledge graphs & connected retrieval

What useful question becomes answerable by following a relationship that a keyword search would miss?

Entity resolution & data quality

Which ambiguous names or duplicate records would most distort our decisions if we matched them incorrectly?

Agents & delegated work

What can this agent do independently, what must it escalate, and what evidence shows it completed the job?

Human trust, oversight & decision support

At the point of consequence, can the responsible person understand, refuse or correct the AI recommendation?

Evaluation, observability & AI operations

Which representative cases must keep passing when we change the model, prompt or knowledge source?

Operational AI governance

Can we name every material AI use case, its owner, its purpose and the next review decision?

Agent security & resilience

If an agent follows the wrong instruction, which actions can the surrounding system prevent or reverse?

Data sovereignty, control & access

Beyond the hosting location, who can access the data and what happens if we must change provider?

Adoption, diffusion & organisational change

Which workflow, owner and measure of success will change if this pilot becomes ordinary work?

AI across software delivery

Where is the real delivery bottleneck, and what evidence will show that AI improves the whole flow?

AI in healthcare & clinical work

What specific task is being proposed, and what kind of evidence would show that it helps in the intended setting?

Creative authorship & human expression

Which creative decisions remain with people, and how can an audience tell whose voice they are encountering?

Conversational voice & multimodal interaction

What should happen when a person interrupts, changes their mind or asks for evidence while a tool is running?

AI in the physical world

Which physical constraint or failure mode would invalidate a convincing software-only demonstration?

Knowledge readiness & authority

The work of making organisational knowledge discoverable, understandable, current and attributable to someone who can stand behind it.

Why it matters

An AI can retrieve a document without knowing whether it is the approved version. Useful organisational AI needs decisions about ownership, access, conflicting sources and updates as well as a capable model.

Keep these distinctions clear

- Information availability is different from authority: a readable file may be obsolete or unapproved.
- A citation shows where a statement came from; it does not by itself establish that the statement is true.
- A larger corpus improves coverage only when relevant material can be found and interpreted.

A question to take into your work

For one important business question, which source wins when two documents disagree, and who decides?

Where it connects at AICON

AI Is Ready. Is Your Knowledge? Why the models we fixate on are rarely the real constraint in creating value from AI — [programme topic link](#)

The Importance of Data and Governance as the Foundations for AI Success — [programme topic link](#)

Follow the idea

supplies trusted context to → Retrieval-augmented generation: Retrieval needs usable source knowledge, including its scope and authority.

requires ownership from → Operational AI governance: Someone must own source approval, refresh and access decisions.

Relationships above are editorial synthesis for learning.

Knowledge readiness & authority

Selected reading and watching

Liberty IT: our journey to responsible generative AI

vendor case study

See how reusable capabilities, use cases and governance were developed together.

Provider-reported experience; outcomes have not been independently verified here.

Tom Swann: observable data quality with Elementary and DataHub

practitioner article

Connect quality checks with discoverable data context and ownership.

Practitioner explanation or experience, not an independent controlled study.

Instil: a practical guide to RAG strategies

practitioner article

Understand how source material enters an answer through retrieval.

Practitioner explanation or experience, not an independent controlled study.

Stanford MS&E435 Economics of the AI Supercycle | Spring 2026 | Enterprise Internal Knowledge

video · Stanford Online · 48 min · 2026-05-22

Connect the conference knowledge demonstration to the challenge of making internal business knowledge useful to AI.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Effective context engineering for AI agents

practitioner article · Anthropic · 2025-09-29

Turn the knowledge-constraint argument into practical design choices.

Public article reviewed. Background learning; no AICON appearance or endorsement implied.

Retrieval-augmented generation

Finding relevant material and supplying it to a language model as context for an answer, instead of expecting the model to remember everything.

Why it matters

RAG is a practical bridge between a model and organisational documents. Its quality depends on what is indexed, how documents are divided, how candidates are ranked and whether the final answer uses the evidence correctly.

Keep these distinctions clear

- Retrieval changes the context supplied at use time; fine-tuning changes model parameters.
- Similarity is a useful search signal, not a guarantee of relevance or truth.
- Reranking can improve candidate selection but adds work and latency.

A question to take into your work

Can the system retrieve the passage that would let a knowledgeable colleague answer this question?

Where it connects at AICON

Zero ontology — programme topic link

AI Is Ready. Is Your Knowledge? Why the models we fixate on are rarely the real constraint in creating value from AI — editorial thematic connection

Follow the idea

can be extended by → Knowledge graphs & connected retrieval: Relationships can supply context that isolated passages omit.

must be tested through → Evaluation, observability & AI operations: Assess retrieval coverage separately from answer correctness.

Relationships above are editorial synthesis for learning.

Retrieval-augmented generation

Selected reading and watching

Instil: practical RAG strategies

practitioner article

Compare basic retrieval, reranking and more elaborate retrieval patterns.

Practitioner explanation or experience, not an independent controlled study.

Instil: why chunking matters

practitioner article

Compare fixed, recursive, semantic and late chunking approaches.

Practitioner explanation or experience, not an independent controlled study.

GraphRAG local search

technical documentation

See how entities, relationships and source text can be combined for a specific question.

Describes this implementation, not a comparative evaluation.

Stanford CS25: V3 I Retrieval Augmented Language Models

video · Stanford Online · 79 min · 2024-01-25

Understand why retrieving relevant organisational knowledge is different from relying on model training.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Retrieval-Augmented Generation for Knowledge-Intensive NLP Tasks

research paper · Patrick Lewis and co-authors / arXiv · 2020-05-22

Understand the original architecture behind answers grounded in retrieved material.

Abstract reviewed; linked full paper not newly reviewed in this supplement. Background learning; no AICON appearance or endorsement implied.

Knowledge graphs & connected retrieval

A representation of entities and explicit relationships that lets people or software follow how pieces of knowledge connect.

Why it matters

A graph can connect a topic to a session, a speaker to their work, or a claim to its evidence. Its value comes from the meaning and quality of those connections. A visually impressive network alone is not a useful knowledge model.

Keep these distinctions clear

- A graph visualisation displays relationships; a graph retrieval system uses them to select context.
- Local questions concern particular entities; global questions ask for themes across a collection.
- An editor-curated thematic connection is different from a relationship explicitly asserted in a source.

A question to take into your work

What useful question becomes answerable by following a relationship that a keyword search would miss?

Where it connects at AICON

Zero ontology — programme topic link

AI Is Ready. Is Your Knowledge? Why the models we fixate on are rarely the real constraint in creating value from AI — editorial thematic connection

Follow the idea

depends on consistent identities from → Entity resolution & data quality: Incorrectly merged or duplicated entities distort the graph.

makes relationships visible within → Knowledge readiness & authority: Connections can expose ownership, provenance and missing context.

Relationships above are editorial synthesis for learning.

Knowledge graphs & connected retrieval

Selected reading and watching

From Local to Global: a GraphRAG approach to query-focused summarization

research paper

Read the graph and community-summary approach to questions about whole collections.

The captured full paper is a preprint; its evaluated global-question results are not a blanket advantage over all RAG systems.

GraphRAG global search

technical documentation

Understand how community reports support collection-wide answers.

Describes this implementation, not a comparative evaluation.

GraphRAG DRIFT search

technical documentation

Explore an approach that combines community context with more detailed retrieval.

Describes this implementation, not a comparative evaluation.

Tom Swann: exploring AWS Neptune

practitioner article

See a practitioner introduction to graph database ideas.

Practitioner explanation or experience, not an independent controlled study.

Intro to GraphRAG

video · Microsoft Reactor · 61 min · 2024-09-07

See how relationships can support retrieval; pair with the research paper to understand the evidence and limitations.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Entity resolution & data quality

Deciding when two records refer to the same real-world thing, while preserving differences that matter.

Why it matters

A person, organisation or product can appear under several names. Joining those records can reveal connections; joining the wrong records manufactures connections. This is a central quality problem for any useful graph.

Keep these distinctions clear

- Matching names is not the same as identifying the same entity.
- A false merge invents a relationship; a missed merge hides one.
- A model-generated match should carry evidence and uncertainty rather than becoming an unquestioned fact.

A question to take into your work

Which ambiguous names or duplicate records would most distort our decisions if we matched them incorrectly?

Where it connects at AICON

Zero ontology — programme topic link

Follow the idea

establishes identities for → Knowledge graphs & connected retrieval: Edges need stable endpoints to mean what they appear to mean.

needs error analysis from → Evaluation, observability & AI operations: Test false merges and missed matches separately.

Relationships above are editorial synthesis for learning.

Entity resolution & data quality

Selected reading and watching

Deep Entity Matching with Pre-Trained Language Models

research paper

Study Ditto and the use of language models for matching record pairs.

Research findings apply to the studied setting; this is not a universal performance guarantee.

Entity Matching using Large Language Models

research paper

Explore LLM matching and the effects of prompts and explanations.

Research findings apply to the studied setting; this is not a universal performance guarantee.

Observable data quality with Elementary and DataHub

practitioner article

Place individual matching decisions in the wider discipline of monitored data quality.

Practitioner explanation or experience, not an independent controlled study.

Agents & delegated work

AI systems that pursue a goal through a sequence of decisions and tool actions, sometimes coordinating with other agents.

Why it matters

Delegation can connect reasoning to real work, but every added action creates another place where an error can propagate. Designing the workflow includes defining tools, authority, stopping conditions and handover points.

Keep these distinctions clear

- A chatbot response and an action in a business system have different consequences.
- Multiple agents add coordination choices; they do not automatically improve results.
- A technically available tool is not the same as permission to use it for every purpose.

A question to take into your work

What can this agent do independently, what must it escalate, and what evidence shows it completed the job?

Where it connects at AICON

[The Agent at Work: Trust and Capability in the Age of Delegation — programme topic link](#)

[From AI Assistants to Agentic Delivery — programme topic link](#)

[Autonomy at the Frontier — programme topic link](#)

Follow the idea

creates new oversight demands for → Human trust, oversight & decision support: Delegation changes what users must understand and supervise.

expands the action surface governed by → Agent security & resilience: Tool permissions and connected systems shape the consequences of errors.

Relationships above are editorial synthesis for learning.

Agents & delegated work

Selected reading and watching

Instil: making sense of multi-agent architectures

practitioner article

Compare network, supervisor and hierarchical coordination trade-offs.

Practitioner explanation or experience, not an independent controlled study.

Kainos: the promise of AI agents

practitioner article

Explore proposed applications and delivery considerations from a provider perspective.

Practitioner explanation or experience, not an independent controlled study.

Enzai: agentic AI governance

vendor case study

Examine a concrete product framing of autonomy, allowed actions and escalation.

Provider-reported experience; outcomes have not been independently verified here.

Tips for building AI agents

video · Anthropic · 18 min · 2025-02-13

Decide where autonomous tool use earns its additional complexity, cost and oversight.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Context Engineering for AI Agents with LangChain and Manus

video · LangChain · 61 min · 2025-10-14

Learn practical ways to manage limited context instead of putting every document into every request.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Building effective agents

practitioner article · Anthropic · 2024-12-19

Choose the simplest architecture that improves measured outcomes.

Public article reviewed. Page notes that tooling has changed since original publication; use patterns as background, not frozen implementation instructions. Background learning; no AICON appearance or endorsement implied.

Human trust, oversight & decision support

Designing AI-assisted work so people understand the system's role, can challenge it and can intervene meaningfully.

Why it matters

A human approval button is weak oversight if the reviewer lacks time, context or the ability to stop an action. Trust becomes useful when it is calibrated to what the system can actually do and the evidence available.

Keep these distinctions clear

- Confidence in an interface is different from demonstrated reliability.
- Human involvement is different from a human having meaningful control.
- Decision support and delegated decision-making require different responsibilities.

A question to take into your work

At the point of consequence, can the responsible person understand, refuse or correct the AI recommendation?

Where it connects at AICON

Why agents make designing for trust harder, yet even more critical — [programme topic link](#)

What Artists Can Teach Us About Building AI People Will Actually Trust? — [programme topic link](#)

Hercules — [programme topic link](#)

Regulating Intelligence — [programme topic link](#)

Follow the idea

should be calibrated by → Evaluation, observability & AI operations: Measured performance helps people decide when to rely on a system.

needs accountability defined by → Operational AI governance: Oversight must have a named owner and a usable escalation route.

Relationships above are editorial synthesis for learning.

Human trust, oversight & decision support

Selected reading and watching

NIST AI Risk Management Framework 1.0

framework

Use the framework to connect context, measurement and management of AI risks.

A framework for structuring work, not proof of compliance or effectiveness.

NIST Generative AI Profile

framework

Explore risks and suggested actions specific to generative AI.

A framework for structuring work, not proof of compliance or effectiveness.

AI and the future of psychiatry: qualitative findings

research paper

Read how clinicians discussed empathy, collaboration and the impact of AI.

This analyses responses from a 2019 physician survey; attitudes are not evidence of current model capability or patient outcomes.

How to Construct Domain Specific LLM Evaluation Systems: Hamel Husain and Emil Sedgh

video · AI Engineer · 19 min · 2024-09-19

Translate an abstract quality score into tests that reflect a particular product and its users.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Evaluation, observability & AI operations

Testing whether an AI system does the intended job, and recording enough of its real operation to understand changes and failures.

Why it matters

Model, prompt and source changes can alter behaviour even when the application still runs. Teams need examples of acceptable performance, traceable versions and a way to detect regressions after release.

Keep these distinctions clear

- Observability records what happened; evaluation judges it against a purpose or criterion.
- Offline tests and production monitoring answer different questions.
- An LLM judge is another fallible measurement method, not ground truth.

A question to take into your work

Which representative cases must keep passing when we change the model, prompt or knowledge source?

Where it connects at AICON

AI Adoption To AI Operations — [programme topic link](#)

Building a Real-Time Voice Agent that Feels Conversational, Not Transactional — [programme topic link](#)

Follow the idea

provides operating confidence for → Adoption, diffusion & organisational change: Production use needs repeatable evidence, not a successful demonstration.

diagnoses failures in → Retrieval-augmented generation: Separate missing evidence from unsupported answer generation.

Relationships above are editorial synthesis for learning.

Evaluation, observability & AI operations

Selected reading and watching

Tom Swann: AI Adoption to AI Operations

practitioner article

Read a conference contributor's account of moving from integrations to measurement and operational reliability.

Practitioner explanation or experience, not an independent controlled study.

Langfuse: evaluation of LLM applications

technical documentation

Compare evaluation workflows and the roles of datasets and evaluators.

Describes this implementation, not a comparative evaluation.

Langfuse: LLM observability and application tracing

technical documentation

Understand the trace information needed to investigate an AI workflow.

Describes this implementation, not a comparative evaluation.

Langfuse: link prompts to traces

technical documentation

Connect observed behaviour with the prompt version that produced it.

Describes this implementation, not a comparative evaluation.

How To Approach Your AI Evals

video · Hamel Husain · 4 min · 2026-06-16

A quick entry point before the longer domain-specific evaluation talk; the description primarily provides presenter background.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

Demystifying evals for AI agents

practitioner article · Anthropic · 2026-01-09

Test the entire agent interaction rather than treating one plausible answer as proof of quality.

Public article reviewed. Background learning; no AICON appearance or endorsement implied.

Operational AI governance

The people, decisions and routines that determine how AI is introduced, used, monitored and changed.

Why it matters

A policy document cannot on its own keep track of deployed systems or resolve a difficult use case. Governance becomes operational through inventories, accountable owners, risk assessment, review and follow-through.

Keep these distinctions clear

- An inventory tells you what exists; an assessment examines its risks and context.
- A voluntary framework is different from a legal obligation.
- A vendor's compliance feature is not evidence that a customer's implementation complies.

A question to take into your work

Can we name every material AI use case, its owner, its purpose and the next review decision?

Where it connects at AICON

The Quiet Risk: Why AI Governance Does Not End With the Policy — [programme topic link](#)

Building an Enterprise AI Governance Programme in Practice — [programme topic link](#)

The Importance of Data and Governance as the Foundations for AI Success — [programme topic link](#)

Follow the idea

defines delegation limits for → Agents & delegated work: Agent actions need owners and decision boundaries.

assigns responsibility for → Knowledge readiness & authority: Source ownership is part of operational control.

Relationships above are editorial synthesis for learning.

Operational AI governance

Selected reading and watching

NIST AI Risk Management Framework 1.0

framework

Use Govern, Map, Measure and Manage to structure the work.

A framework for structuring work, not proof of compliance or effectiveness.

NIST AI RMF Playbook

framework

Move from framework headings to practical suggested actions.

A framework for structuring work, not proof of compliance or effectiveness.

Enzai: how to build an AI system inventory

practitioner article

Consider the data and ownership a working inventory needs.

Practitioner explanation or experience, not an independent controlled study.

Liberty IT: responsible generative AI journey

vendor case study

See an organisation's account of embedding governance in development.

Provider-reported experience; outcomes have not been independently verified here.

Agent security & resilience

Protecting the information, tools and operational boundaries an AI system depends on, and preparing to contain failures.

Why it matters

An agent's risk depends on the systems it can reach and the authority it has. A useful security discussion follows the complete workflow, including external content, tool execution, logs and recovery.

Keep these distinctions clear

- A model refusing a request is different from a system enforcing a permission boundary.
- Privacy, cybersecurity and service availability overlap but are not interchangeable.
- A product security claim is not an independent assurance assessment.

A question to take into your work

If an agent follows the wrong instruction, which actions can the surrounding system prevent or reverse?

Where it connects at AICON

Autonomy at the Frontier — [programme topic link](#)

The Agent at Work: Trust and Capability in the Age of Delegation — [programme topic link](#)

Follow the idea

constrains tool use by → Agents & delegated work: Permissions limit what delegated work can change.

overlaps with control questions in → Data sovereignty, control & access: Both require examining access and dependencies, not just model behaviour.

Relationships above are editorial synthesis for learning.

Agent security & resilience

Selected reading and watching

NIST Generative AI Profile

framework

Review generative-AI risk categories and proposed mitigation actions.

A framework for structuring work, not proof of compliance or effectiveness.

Enzai: agentic AI governance

vendor case study

Inspect a provider's approach to allowed actions and escalation.

Provider-reported experience; outcomes have not been independently verified here.

Options: PrivateMind

vendor case study

Compare claims about private deployment, confidential computing and observability.

This is a product description; no penetration test or independent verification of its security claims is included.

Data sovereignty, control & access

Understanding who controls data and AI infrastructure, who can access them and which dependencies shape that control.

Why it matters

A data-centre location is only one part of an AI deployment. Operational access, provider dependencies and the ability to move or run a workload also affect the control an organisation retains.

Keep these distinctions clear

- Residency describes where data is stored; sovereignty asks wider questions about control and exposure.
- Open weights do not by themselves establish ownership of training data or remove operating responsibilities.
- Private infrastructure still requires access management and operational assurance.

A question to take into your work

Beyond the hosting location, who can access the data and what happens if we must change provider?

Where it connects at AICON

Residency Isn't Sovereignty: Own Your Exposure — [programme topic link](#)

The AI Capability Divide — [programme topic link](#)

Story Engine: AI for a Civic Heritage Destination at Scale — [programme topic link](#)

Follow the idea

requires controls from → Agent security & resilience: Control claims need an account of actual access and operational protections.

shapes deployment choices in → Adoption, diffusion & organisational change: Infrastructure and access constraints influence which uses are feasible.

Relationships above are editorial synthesis for learning.

Data sovereignty, control & access

Selected reading and watching

Options: PrivateMind

vendor case study

Use a concrete private-deployment proposition to formulate questions about control.

Provider-reported experience; outcomes have not been independently verified here.

AI for NI: a strategic overview

policy report

Read the infrastructure, data and skills context behind regional capability.

Contextual reading; not proof of a particular product outcome.

Open AI meets open notes

research paper

Examine how access to information can interact with privacy risks.

An ethics analysis concerning patient records; it is not a technical assessment of a particular deployment.

Adoption, diffusion & organisational change

Turning AI capability into sustained use across teams and organisations, with the skills, workflows and ownership needed to make it useful.

Why it matters

A promising pilot can stall when it reaches existing processes. Diffusion asks how benefits spread beyond an expert team; adoption asks what changes in day-to-day work.

Keep these distinctions clear

- Tool access is different from effective use.
- A faster task is different from an improved end-to-end business outcome.
- An organisation's case study is context, not a promised return for another organisation.

A question to take into your work

Which workflow, owner and measure of success will change if this pilot becomes ordinary work?

Where it connects at AICON

Diffusion Wins — [programme topic link](#)

The Adoption Gap — [programme topic link](#)

AI Readiness Workshop — [programme topic link](#)

Follow the idea

is constrained by → Knowledge readiness & authority: Teams need usable business knowledge as well as tools.

needs evidence from → Evaluation, observability & AI operations: Measures help distinguish sustained value from initial enthusiasm.

Relationships above are editorial synthesis for learning.

Adoption, diffusion & organisational change

Selected reading and watching

AI for NI: a strategic overview

policy report

Explore the regional combination of data, trust, infrastructure and skills.

Contextual reading; not proof of a particular product outcome.

Liberty IT: responsible generative AI journey

vendor case study

Read an account of experimentation becoming reusable organisational capabilities.

Provider-reported experience; outcomes have not been independently verified here.

Tom Swann: AI Adoption to AI Operations

practitioner article

Follow the operational work that appears after initial AI features ship.

Practitioner explanation or experience, not an independent controlled study.

Stanford MS&E435 Economics of the AI Supercycle | Spring 2026 | Enterprise Internal Knowledge

video · Stanford Online · 48 min · 2026-05-22

Connect the conference knowledge demonstration to the challenge of making internal business knowledge useful to AI.

Publisher title and description only; video not watched and transcript not reviewed. Background learning; no AICON appearance or endorsement implied.

AI across software delivery

Using AI in the chain from requirements and design through implementation, testing and operation.

Why it matters

Code generation is one part of delivery. Changes to requirements, review, system design and quality control determine whether faster local work becomes a better delivered product.

Keep these distinctions clear

- Producing more code is different from delivering a useful change.
- An assistant helping a developer is different from an agent coordinating a delivery workflow.
- A legacy-code explanation is not evidence that a proposed migration preserves behaviour.

A question to take into your work

Where is the real delivery bottleneck, and what evidence will show that AI improves the whole flow?

Where it connects at AICON

Beyond Coding — [programme topic link](#)

From AI Assistants to Agentic Delivery — [programme topic link](#)

Follow the idea

can delegate parts of delivery through → Agents & delegated work: Tool-using systems extend AI beyond drafting code.

needs regression evidence from → Evaluation, observability & AI operations: Generated changes must preserve required behaviour.

Relationships above are editorial synthesis for learning.

AI across software delivery

Selected reading and watching

Kainos: AI Augmented Delivery

practitioner article

Read a provider's argument for connecting speed to trusted delivery outcomes.

Practitioner explanation or experience, not an independent controlled study.

Instil: LLM-assisted code modernisation

practitioner article

Explore the use of language models in understanding and modernising legacy code.

Practitioner explanation or experience, not an independent controlled study.

Instil: agentic software delivery lifecycle

vendor case study

Inspect a proposed end-to-end delivery approach rather than a single coding tool.

Provider-reported experience; outcomes have not been independently verified here.

AI in healthcare & clinical work

Applying AI to healthcare tasks while examining the evidence, responsibilities and patient context of each particular use.

Why it matters

Healthcare discussions often jump from capability demonstrations to claims about replacing clinicians. Useful analysis separates tasks, professional attitudes, actual clinical performance and patient outcomes.

Keep these distinctions clear

- A survey of clinicians measures opinions, not model accuracy.
- Administrative assistance and clinical decision support have different purposes and consequences.
- An abstract, preprint and journal version of the same study are not three independent studies.

A question to take into your work

What specific task is being proposed, and what kind of evidence would show that it helps in the intended setting?

Where it connects at AICON

Dr Bot: Why AI Is Coming for Doctors' Jobs? — [programme topic link](#)

AI in Health: What Getting It Right Actually Looks Like — [programme topic link](#)

Follow the idea

makes the consequences concrete for → Human trust, oversight & decision support: Clinical work exposes the difference between help, reliance and responsibility.

requires context-specific assessment through → Operational AI governance: Intended use and affected people shape the review needed.

Relationships above are editorial synthesis for learning.

AI in healthcare & clinical work

Selected reading and watching

AI and the future of psychiatry: insights from a global physician survey

research paper

Explore how psychiatrists expected different tasks to change.

Survey of professional opinions collected in 2019; not a clinical effectiveness study or a measure of present AI capability.

AI and the future of psychiatry: qualitative findings

research paper

Read the themes in clinicians' written responses about empathy and collaboration.

A qualitative analysis of the same global survey programme, not independent replication of clinical outcomes.

Open AI meets open notes

research paper

Examine the intersection of patient access, online resources and privacy.

Ethical analysis rather than a clinical trial.

Creative authorship & human expression

Using generative tools within creative work while keeping sight of who sets the intent, makes choices and speaks through the result.

Why it matters

Creative AI raises questions beyond output quality: authorship, participation, provenance and whose perspective is represented. The conference connects this to both production workflows and trust.

Keep these distinctions clear

- Rendering an output is different from setting the creative intent.
- A generated reconstruction is different from a preserved first-person account.
- A provider's account of a creative project is not an independent audience-impact study.

A question to take into your work

Which creative decisions remain with people, and how can an audience tell whose voice they are encountering?

Where it connects at AICON

Generative AI is the Renderer Not the Author — [programme topic link](#)

From AI Production to AI Platform — [programme topic link](#)

What Artists Can Teach Us About Building AI People Will Actually Trust? — [programme topic link](#)

Follow the idea

reveals design questions in → Human trust, oversight & decision support: People need to understand their agency and the system's role.

needs provenance from → Knowledge readiness & authority: Creative source material and human contributions need attributable context.

Relationships above are editorial synthesis for learning.

Creative authorship & human expression

Selected reading and watching

Hamilton Robson: Listening to Belfast

vendor case study

Examine an AI interviewing project's stated intention to preserve participants' own stories.

Provider-reported experience; outcomes have not been independently verified here.

Hamilton Robson: Powering Belfast's Human Narrative with AI

vendor case study

Explore the relationship between technical delivery and human narrative.

Provider-reported experience; outcomes have not been independently verified here.

NIST Generative AI Profile

framework

Place creative use within a wider discussion of generative-AI risks and human context.

A framework for structuring work, not proof of compliance or effectiveness.

Conversational voice & multimodal interaction

AI interaction that combines speech, turn-taking and sometimes other inputs with retrieval or tool use.

Why it matters

A conversational voice experience is a workflow, not just a natural-sounding voice. Interruptions, timing, tool responses and evidence all affect whether it helps the user complete the task.

Keep these distinctions clear

- Speech recognition, speech generation and speech-to-speech interaction are different components or approaches.
- A fluent voice is not evidence of a correct answer.
- A product guide explains implementation; it does not establish superiority over alternatives.

A question to take into your work

What should happen when a person interrupts, changes their mind or asks for evidence while a tool is running?

Where it connects at AICON

Building a Real-Time Voice Agent that Feels Conversational, Not Transactional — [programme topic link](#)

Follow the idea

provides an interaction channel for → Agents & delegated work: Voice requests can initiate tool-using workflows.

needs interaction tests from → Evaluation, observability & AI operations: Correctness alone misses timing, interruption and recovery failures.

Relationships above are editorial synthesis for learning.

Conversational voice & multimodal interaction

Selected reading and watching

Amazon Nova Sonic: core concepts

technical documentation

Understand bidirectional streaming and the interaction model.

Describes this implementation, not a comparative evaluation.

Amazon Nova: voice conversation prompts

technical documentation

Explore how instructions shape spoken interaction.

Describes this implementation, not a comparative evaluation.

AWS: Live Meeting Assistant with Transcribe, Bedrock and Strands

technical documentation

Read a worked architecture connecting live speech with an agent workflow.

Describes this implementation, not a comparative evaluation.

AI in the physical world

AI systems whose perception or decisions interact with equipment, manufacturing processes or physical environments.

Why it matters

Physical applications connect software behaviour with sensors, hardware and operating conditions. A useful discussion must ask how an idea survives real constraints and how failure is detected and contained.

Keep these distinctions clear

- A digital demonstration is different from operation in a changing physical environment.
- Perception, planning and physical action are distinct parts of a system.
- General AI risk guidance does not substitute for domain-specific engineering validation.

A question to take into your work

Which physical constraint or failure mode would invalidate a convincing software-only demonstration?

Where it connects at AICON

AI in the Physical World: Prompts are Cheap, Parts Aren't — programme topic link

Follow the idea

needs environment-specific testing from → Evaluation, observability & AI operations: Software tests alone cannot establish performance in a physical setting.

needs responsibilities from → Operational AI governance: Operational ownership and intervention arrangements matter where actions have physical effects.

Relationships above are editorial synthesis for learning.

AI in the physical world

Selected reading and watching

NIST AI Risk Management Framework 1.0

framework

Start with intended context, measurement and risk management.

A framework for structuring work, not proof of compliance or effectiveness.

RT-2: Vision-Language-Action Models Transfer Web Knowledge to Robotic Control

research paper · Brohan and co-authors / arXiv · 2023-07-28

See how vision-language pretraining is combined with robot trajectories and action tokens to connect language instructions to physical control.

Abstract and official project explanation reviewed; reported research settings do not establish industrial safety or general deployment readiness. Background learning; no AICON endorsement implied.

Open X-Embodiment: Robotic Learning Datasets and RT-X Models

research project · Open X-Embodiment Collaboration

Explore how standardised robot datasets and shared policies support transfer across different robot platforms.

Official project methods/results text reviewed; experimental transfer is not proof of dependable operation in every environment. Embedded demonstrations not watched. Background learning; no AICON endorsement implied.