

EU AI Act, Article 14 Readiness Checklist for Contact Centers

EU AI Act, Article 14 sets out five specific oversight capabilities. The checklist below covers those five, along with what it actually takes to put them into practice, day-to-day operations, vendor evaluation, and documentation. This is not an extensive checklist but helps organizations to consider their readiness.

This checklist accompanies our blog, [What the EU AI Act's Article 14 Demands of Human + AI Systems](#).

1. Understand your Exposure

- ☐ Do any of your customer-facing AI systems qualify as high-risk under the Act? (see Annex III for the specific high-risk categories, and Article 6 for how that classification works)
- ☐ Do these AI systems touch EU customers/operations, even if your business is US-based?
- ☐ Is your contact center acting as a deployer of these AI systems, a provider, or both?
- ☐ Which obligations sit with you versus your AI vendor?
- ☐ Does your AI vendor understand the implications of the EU AI Act?
- ☐ Have you made any customizations to a vendor's AI or oversight tooling? (ie. if you've changed how oversight functions rather than just how it's configured, treat that as a design decision, not a settings change.)
- ☐ Have you developed training material or courses for human advisors in the contact center that reflect the EU AI Act, Article 14?

2. Build Real Understanding

- ☐ Do the people you have assigned for AI oversight actually understand the AI system's capabilities and known failure modes?
- ☐ Have you defined what anomalous or unexpected behavior actually looks like for your specific use cases, and do your human overseers know how to recognize it?
- ☐ Are your human overseers trained specifically on automation/AI bias?
- ☐ Is this understanding refreshed when the AI system is updated, or is it based on how the system behaved when oversight training first happened?

3. Make Interpretation Possible

- ☐ Have your human overseers been trained on the use of confidence scores, flagged intents, or other AI-generated triggers?
- ☐ Does the AI system itself provide the tools needed to interpret its output, like a confidence score or reasoning context, or does interpretation depend entirely on what your team can infer without that support?
- ☐ Can overseers correctly interpret ambiguous or edge-case output?
- ☐ Do your human overseers know how to react if they are unsure how to interpret an AI or edge-case output?

4. Verify Override and Intervention

- ☐ Can a human overseer disregard, override, or reverse an AI's output on a case-by-case basis?
- ☐ Does a stop mechanism exist that lets a human halt the AI mid-interaction, not just remove it from future ones?
- ☐ Have you tested override and halt capabilities under real operating conditions, including at volume, not just in a demo or sandbox environment?

5. Choose and Document your Oversight Model

- ☐ Where is full takeover required?
- ☐ Where is supervisory review appropriate, based on the risk level of the interaction?
- ☐ For supervisory review, can the human monitor, approve, or edit AI output without losing the efficiency of keeping the AI in the loop?
- ☐ Is full takeover available as a fallback even where supervisory review is the default?
- ☐ Document the reasoning behind which model applies where. This also becomes part of your compliance record and training guide.

6. Close the Design-Risk Gap

- ☐ If using an AI vendor, is oversight capability built into the system's architecture?
- ☐ Does the AI vendor satisfy each of the 5 requirements of Article 14 specifically and not just a general oversight statement?
- ☐ Where a hybrid agentic approach is in place, AI handling routine work with a human able to step in at the point of interaction, review, or output, have the handoff points, monitoring visibility, and override mechanisms been designed in from the start?

7. Operationalize

- ☐ Have you run end-to-end tests of human oversight, override, and halt capabilities under real interaction conditions, not just in isolation?
- ☐ Have you confirmed oversight responsibilities are assigned to specific, trained individuals?
- ☐ Are you keeping records of oversight design decisions, vendor documentation, and testing results, since this becomes your evidence of compliance if it's ever questioned?

For more information about how ServisBOT addresses
Human + AI Oversight and Collaboration
please [contact us](#)

This is the second piece in our Hybrid Agentic AI series. Read the first, [The Hybrid Agentic AI Model: How Human and AI Collaboration Is Changing the Way Regulated Businesses Deploy AI](#), for more on where this approach started.

