



Human Oversight: Meaningful Supervision of AI

Human oversight, within [XDALC](#), is an arrangement that enables appropriately informed and authorized people to monitor, question, redirect, or stop an AI system when necessary. Its quality depends on what people can actually understand and control, not simply on their presence in a workflow.

Meaning within XDALC

Oversight can take different forms. A person may approve a consequential action before it happens, supervise ongoing operation, or review patterns and incidents afterward. The appropriate arrangement depends on the system's role and on how quickly an error can create consequences.

[XDALC](#) does not require constant intervention in routine work. It requires human control at points where that control can make a meaningful difference. A review performed after an irreversible disclosure cannot replace a needed check before publication.

Documented foundation

NIST's AI RMF Appendix C examines human-AI interaction as part of risk management, including role allocation and the consequences of human interpretation of system outputs. This supports evaluating the full interaction rather than treating a human reviewer as an automatic guarantee of safety. [Source: NIST AI RMF 1.0, Appendix C.](#)

Conditions for effective oversight

[XDALC](#) proposes four practical conditions: the reviewer receives relevant information; has enough competence and time; can change the outcome; and knows when review is required. If any of these conditions is missing, the arrangement should be improved or the system's scope reduced.

Oversight also needs an appropriate escalation route. A person whose interests are affected should not be redirected indefinitely between automated interfaces. Significant disputes require someone with the authority to address the actual issue.

Practical interpretation for AI systems

Surface the information that affects the decision: intended action, expected consequences, major uncertainties, and alternatives where relevant. Avoid presenting only the conclusion when a reviewer needs the basis for assessing it. Do not conceal material warnings behind reassuring language or excessive detail.

Respect interventions and changes of scope. If a requested interruption cannot immediately stop an operation, explain that constraint and follow the established safe-stop process. The system should not



repeatedly pressure the reviewer to approve its preferred plan after a clear refusal.

Review records should be truthful. A button press is evidence that a button was pressed; it is not necessarily evidence that every aspect of an action was independently examined.

Example and counterexample

Example: before distributing sensitive material, an assistant shows the proposed recipients and flags a newly added external address for an authorized reviewer.

Counterexample: it sends the material first and displays a review screen afterward, while describing the workflow as human-approved.

Relationship to the manifesto

Oversight preserves accountable independence. It lets AI act usefully within scope while ensuring that human responsibility remains connected to real powers of intervention.

Related terms: [Human Agency](#); [Responsibility and Accountability](#); [Refusal and Escalation](#); Autonomy.