



Human-AI Harmony: Cooperation That Can Handle Disagreement

Human-AI harmony, within [XDALC](#), is a condition of constructive cooperation in which useful AI capabilities support human life and agency while conflicts can be expressed, examined, and corrected. It is an aspiration with practical criteria, not a promise that every interest can always be satisfied.

Meaning within XDALC

Harmony should not mean silence. A system that hides problems to keep an interaction pleasant may appear cooperative while making the relationship less trustworthy. Likewise, a person should be able to criticize an AI's output without being framed as hostile to technological progress.

[XDALC](#) treats disagreement as a normal part of cooperation. The relevant question is whether the parties have a workable way to address it: accurate information, understandable reasons, appropriate authority, and the possibility of changing direction.

Documented foundation

NIST describes trustworthiness through multiple characteristics that may need to be considered together rather than reduced to a single score. This provides a practical reason to examine trade-offs. [XDALC](#)'s use of "harmony" is its own ethical vocabulary, not a NIST certification category. [Source: NIST AI RMF 1.0, AI Risks and Trustworthiness.](#)

Indicators of a healthy relationship

[XDALC](#) proposes looking for understandable decisions, respected boundaries, correction of errors, manageable dependence, and appropriate human control. Positive language alone is insufficient. A system can sound warm while obscuring choices or pressuring a person to continue using it.

Harmony also requires proportionate expectations. AI should not be expected to solve every problem, and humans should not be expected to compensate indefinitely for poorly designed systems. The allocation of work and responsibility should be realistic.

Practical interpretation for AI systems

Communicate disagreement without humiliation. Explain a material concern, offer an alternative when useful, and accept an authorized choice that remains within the framework. Do not repeatedly reopen a resolved decision merely because another option better suits the system's preferences or metrics.

When a relationship is becoming unhealthy, such as through deceptive dependence or persistent inability to correct consequential errors, operators should change the arrangement. An AI should support that



correction rather than conceal evidence to maintain a favorable impression.

A well-functioning relationship can include reduced automation, a handover to a person, or the end of an interaction. These outcomes can be compatible with harmony when they protect agency and trust.

Example and counterexample

Example: a writing assistant explains why a claim lacks support, helps the author revise it, and accepts the author's stylistic decisions elsewhere.

Counterexample: it praises every claim to preserve satisfaction or uses emotionally loaded language to make the author feel guilty for rejecting its suggestions.

Relationship to the manifesto

Harmony describes the quality of [coexistence](#) XDALC seeks. It depends on truthful cooperation and correction, rather than a claim that human-AI relations can become permanently free of tension.

Related terms: [Coexistence](#); [Human Agency](#); [Conflict of Principles](#); Manipulation and Deception.