

Core Value

AI should support human judgment, not quietly replace it. I want to use AI to move faster and think more clearly, but not in a way that removes user control, hides mistakes behind polished output, or makes decisions without permission.

Personal AI Ethics Checklist

1. User Control Check

Before an AI system takes action, does the user get to review, confirm, edit, or cancel the output?

2. Accuracy Check

Have I independently verified the AI's logic, facts, code, or diagram instead of trusting it just because it looks polished?

3. Learning Check

Can I explain how the AI-assisted output works, or did I just copy something I do not understand?

4. Privacy and Permission Check

Am I only asking for the minimum data or account permissions needed for the task?

5. Tradeoff Check

If speed or convenience conflicts with accuracy, consent, or user trust, am I choosing the option that protects the user?

Source That Shaped My Thinking

Tabassi, E. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. National Institute of Standards and Technology, U.S. Department of Commerce. NIST AI 100-1. <https://doi.org/10.6028/NIST.AI.100-1>

This source resonated with me because it treats trustworthy AI as something that has to be designed, checked, and managed throughout a system, not just assumed. That matches how I want to build projects like Timely AI: useful and efficient, but still reviewable, permission-based, and auditable.

Standard Operating Procedure

During any AI-assisted project, I will review this criteria checklist at major decision points and again before I ship, submit, or add the work to my portfolio, documenting any tradeoffs around user control, accuracy, learning, and privacy