

Short answer: No, and could be dangerous to try.

Slightly longer answer: With any realistic real-world task assigned to an AGI, there are so many ways in which it could go wrong that trying to block them all off by hand is a hopeless task, especially when something smarter than you is trying to find creative new things to do. You run into the [nearest unblocked strategy](#) problem.

It may be dangerous to try this because if you try and hard-code a large number of things to avoid, it increases the chance that there's a bug in your code that causes major problems, simply by increasing the size of your codebase.

Related

- [☰ Can we constrain a goal-directed AI using specified rules?](#)
- [☰ Wouldn't a superintelligence be smart enough not to make silly mistakes in its ...](#)
- [☰ Can't we just tell an AI to do what we want?](#)
- [☰ Once we notice that a superintelligence is trying to take over the world, can't ...](#)

Alternate phrasings

-