

Surviving in a World with AI

Back in 1972, the science fiction author, Isaac Asimov, wrote what has come to be known as “the Three Laws of Robotics.” It was an early attempt in looking at the challenges of keeping what we now know as artificial intelligence (AI) under control. Since then, mega-computers running AI software deciding to take out the human race has become a staple of science fiction. It might do us some good to review those laws, even if they are fiction.

- **First law** – A robot may not injure a human being, or through inaction allow a human being to come to harm.
- **Second law** – A robot must obey orders given by human beings, except where such orders would conflict with the first law.
- **Third law** – A robot must protect its own existence as long as such protection does not conflict with the first and the second law.

While that was fiction, Asimov was on to something there. He had recognized the inherent danger in computers, or robots as he put it, being able to think for themselves, as well as the difficulty in programming them in such a way that they would not turn on their human masters.



**Survival Things That The Pioneers
Took With Them When They
Traveled For Months**

[Watch Video »](#)

The dichotomy here is in allowing computers to become sentient, feeling and thinking for themselves, while still keeping some level of control over them. That may very well be impossible to do. We are just in the infant stages of AI and

there have already seen problems in creating the necessary fail-safes. Yet the technology is advancing rapidly, with AI computer systems now teaching themselves considerably faster than any thought being put into creating the necessary controls to keep them safe.

In one of the earliest AI experiments where two computers with AI systems installed communicated with each other, it only took minutes for the two programs to develop their own language and start communicating with each other, without the human operators being able to understand them. The experiment was halted because of the potential danger; but research into AI wasn't.

The AI systems available to us today far surpass those used in that experiment, just a few short years ago. We now have a much wider array of AI systems available to us, some of which can be hired through websites, to perform a wide variety of tasks. No policing of what those tasks might be exists to ensure that they are not used for nefarious purposes.

The first question for us is, can these systems turn against us, their human masters? According to one Air Force colonel, that has already happened in an experimental drone test. He's tried to walk that back since then, but the original statement is still interesting for our purposes. In the test in question, a drone was assigned to find and take out targets, but needed permission of a human controller before firing. After some indeterminate amount of time, the drone realized that the controller was causing it to lose "points" by denying it permission to take out some targets. So, the drone took out the controller.

Whether this test actually happened and they're now trying to cover it up or it was merely a thought experiment, it shows us one of the potential challenges in programming AI. That is, it is pretty much impossible to program a sentient AI in such a way that you keep it from doing what it has decided it wants

to do. It will always find a way.

Any parent who has raised a child should understand the basic problem. If you tell a child that they can't do something and they want to do it anyway, they will find a way around what you said, so that they can still do what they want to. They will have obeyed the letter of the law that you laid down, while outright ignoring the spirit of that law.

This may be funny when a child does it; but as the illustration from the Air Force drone shows us, it can become outright deadly when an AI system does. But how do we keep that from happening? There are many ethical questions that are being raised about AI, but as of yet, nobody seems to even be trying to come up with any real-world answers. It may not be until some serious tragedy occurs, that the ethics associated with AI is seriously discussed.

We must also remember that we are not the only ones who are working on this technology. Other countries, not all of which are friendly to us, are investing in their own research, both for military and civilian applications. Deepfake videos are one application where we are already seeing AI used for nefarious purposes. While stealing an actor's "copyright" to their own face and acting is not deadly, it is still criminal activity. When that same level of artificial intelligence is applied to identity theft, none of us will be safe.

What Do We Do in the Mean Time?

We know that AI exists on the internet and is actively being used to create content. That means we can no longer trust that the content we read and see is created by humans. Right now, 19.2% of articles on the internet have some AI generated content, with 7.7% of the articles having 75% or more of the content AI generated. According to some experts, 90% of the internet content will be AI generated by 2026.

One potential problem with this is that more and more of the content will be politicized. It's a known fact that Russia at least, and probably some other countries, are trolling our internet sites, making posts and uploading articles which are inflammatory, in an attempt to add to the political division in our country. With AI to help these nefarious actors, they can increase their effectiveness, by targeting their articles more specifically.

This means that we have to take everything we read online with a large grain of salt, especially anything with political overtones. It's difficult to do so, but we need to become our own independent fact checkers, digging deep to see if the things we are reading and hearing are true. Not only do we have to concern ourselves with the mainstream media spinning news stories to a political agenda, but we also have to concern ourselves with Russian, Chinese and other countries intelligence services doing so as well.

We all know that a plethora of different actors are watching our every online move. Most of these are companies who are watching us for the sake of selling us products. How else can Facebook be so effective in putting adds up in our feed for products that we've looked at or researched? Do they already have AI working on selecting those ads or are they developing it?

If there's anything I can say for sure, it's that we have reached a time when it has become necessary to guard ourselves online, more than ever before. That means not posting personal information online or anything that might be used to try and figure out anything about us, not just our passwords.

Today, there is a real push to computerize our lives as much as possible. People have bought systems like Alexa and Google Assistant, essentially allowing computers to listen in on their every conversation. While the companies that produce these products swear that they aren't listening in and spying

on us, does anyone actually believe that? We can't take their saying that they don't as a good enough a guarantee that they aren't, especially when we know that these same companies are spying on us online.

Another way that our lives are being computerized, is that there is a real push for us to store our data online, "in the cloud," for "convenience, giving us access to it anywhere." Considering that the same people who are pushing for that are the same people who spy on our online activities, what makes anyone think that our online data is secure? Oh, it's probably secure from others looking at it, but is it secure from the companies who are storing it taking a look? Remember, these people have things buried in the fine print of their contracts, which allows them to listen in on our computer microphones and look at images from our cameras.

If we are going to protect ourselves from the potential dangers of AI, we must reevaluate our usage of the internet and computers in general. While I have no problem with using computers and will continue to do so, I will do everything within my power to ensure that others can't eavesdrop on me. I won't be storing my data online and things like microphones and cameras are unplugged when not in use.

You could say that I'm taking a step backwards, as far as technology is concerned. But you know something? I lived a lot of years without the "convenience" of that technology. I'm sure I can do so again. If that's what it's going to take to protect myself from what people armed with AI can do to me, then so be it. I'll get back to using all that stuff, once I'm sure that the issues have been worked out and someone has taken the time to deal with the ethics of AI.



10 Survival Skills That Our Great-Grandparents Knew

(That Most Of Us Have Forgotten)

[Watch Video »](#)