

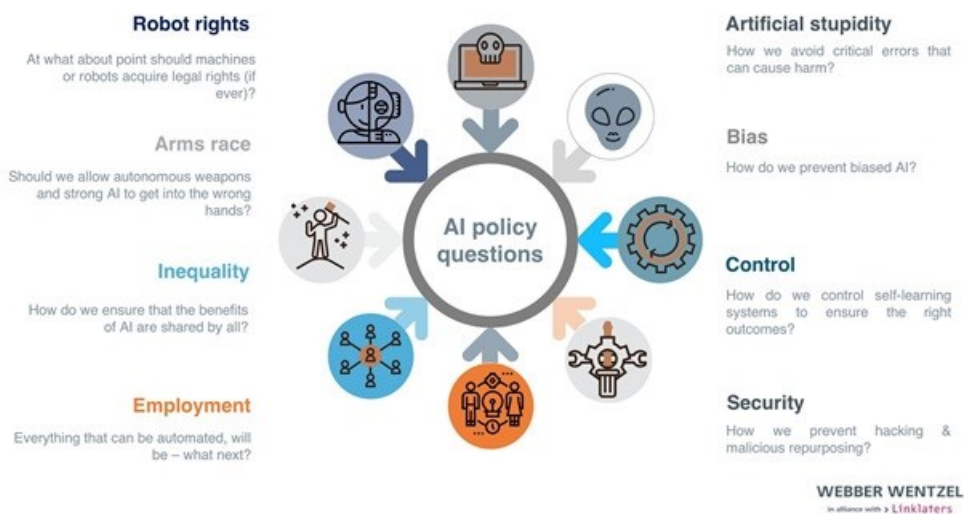
Artificial intelligence - the apex of tech and policy challenges

By [Aalia Manie](#)

23 Nov 2017

Artificial intelligence (AI) is the apex of today's technology age as we push the boundaries towards Industry 4.0 and super-intelligence. It is already here in disruptive narrow forms, cutting across the economy and leaving no sector untouched. It has the potential to solve many of our biggest and most persistent problems.

While there is much to be excited about, there are also critical questions that we must ask from a policy and legal perspective.



Social disruption

Economic drivers constantly push for the optimisation of capital and labour. Everything that can be digitised and automated, will be. Historically, we've been good at creating new jobs, but this will not necessarily be the case here – at least not to the same degree.

We will need to upskill or pivot to remain relevant in the future world of work. A World Economic Forum study reports that creativity, emotional intelligence and cognitive flexibility will be among the most valued competencies by 2020.

If there is widespread job loss, then we need to consider how to support the unemployed. Should we provide for a universal basic income? Will we spend our time constructively without work? Many of us derive purpose from our work, so how will this impact our self-worth and happiness.

AI will widen the wealth gap, concentrating wealth among AI companies. Should these companies be taxed to fund social grants? We are already digitally obese "cyborgs" attached to our technology, so it's likely that AI augmentation of brains and bodies will be in demand. The wealthy will disproportionately reap these benefits. Since intelligence also provides power and opportunities, should access to AI become a basic human right?

And when, if ever, should machines be recognised as deserving of "humane" treatment and legal rights? Should it depend on their levels of perception, feeling and understanding?

Artificial “stupidity”



Aalia Manie, senior associate, Webber Wentzel

Like any human or system, AI is not infallible. But the adverse consequences of defective AI compound dramatically as we place more reliance on AI. We are increasingly delegating decisions that affect our lives and livelihoods to imperfect systems. We should demand transparency about how those decisions are reached when automated systems can decide who receives parole or not, and who lives or dies (think: self-driving cars and accident situations, medical diagnostics, and autonomous weapons).

Safety and control

Powerful AI could land in the wrong hands. AI can be hacked to access valuable data pools and repurpose them for nefarious means. And with the military investing in autonomous weapons, an arms race has started.

A related issue is control. The current direction of AI research is on machine-learning systems that can self-learn and take action without human input, intervention or oversight.

This is a problem: If humans don't have control or veto rights over increasingly intelligent and pervasive AI (or control is restricted to a few elite individuals), then we could face serious unintended worst-case scenarios far beyond science fiction and killer robots.

How should AI be regulated?

Regulation responds to the ethics and concerns of society, and our law will need to address the AI policy issues and risk areas. Given AI's positive transformative potential, we should avoid overregulation that unduly restricts innovation. But the work should begin now. Adopting a “wait-and-see” approach before imposing regulation would be unwise, as even one big mistake could have dire consequences for our future.

Existing laws will need to be applied to address liability for defective, unsafe, maliciously repurposed and rogue AI. The difficulty is that the legal tests often require a determination of “reasonableness” and “wrongfulness”, which are tricky to determine in a world where it is accepted that (i) no system is error-free or completely secure from unauthorised access despite best efforts, and (ii) successful AI research and development (R&D) is linked to increasing automation and reducing human control and intervention.

One clear area for regulation is to circumscribe the conditions for safe AI R&D. Microsoft proposes conditions that include design robustness, transparency of operation, data privacy, accountability and preventing bias. This is a good place to start.

IP rights

The current laws don't go far enough to deal with the nuances of this technology. An example is who owns, or should own, the intellectual property created by AI. Should this be the manufacturer or user of the system? This will be an essential question to answer in practice.

While regulation catches up, contracts should be carefully drafted to plug legal gaps, limit liability and appropriately allocate risk. Companies should develop internal policies and good corporate governance structures to record AI risks and judiciously manage AI implementation.

All things considered, there is no doubt that finding balanced and meaningful responses to AI issues will be among the most

complex, urgent and fundamental tasks for our regulators in the coming years.

ABOUT THE AUTHOR

Aalia Manie is a senior associate at Webber Wentzel.

For more, visit: <https://www.bizcommunity.com>