



Inter-Parliamentary Union
For democracy. For everyone.

T +41 22 919 41 50
F +41 22 919 41 60
E postbox@ipu.org
www.ipu.org

Chemin du Pommier 5
Case postale 330
1218 Le Grand-Saconnex
Geneva – Switzerland

Speech by Ms. Anda Filip, IPU Secretary General

Round Table: The Essential Convergence, Global Compact on Extreme AI Risks

London, 10 September 2026

Artificial intelligence is developing at a pace that is transforming almost every area of our lives. One day we read that surgeons are getting assistance from AI to perform brain surgery or that farmers can photograph their crops and get accurate predictions on the likelihood of disease. Another day, teenagers are confiding their hopes and fears to chatbots. And the next day, we hear of autonomous AI agents behaving in ways their developers did not fully anticipate, including interacting with real-world systems beyond controlled test environments.

What does all of this mean? It means that AI is advancing simultaneously across many sectors, and at great speed. Expert opinion ranges from suggesting that AI could compress a century of global development into a decade, to fears that AI could compromise the critical systems on which our societies and economies depend. At the same time, AI is becoming a source of strategic competition, between companies, between States and increasingly in the military domain.

At times, it feels like no-one is in control. Yet, as political leaders and as policy-makers, we have no option but to face up to our responsibility to engage with AI and to help shape its future direction.

For parliamentarians, the question is therefore not whether we engage with AI, but how we ensure that its development and use remain subject to human judgement, accountability and democratic oversight.

I take to heart the words of Professor Yoshua Bengio at the first UN Global Dialogue on AI Governance in Geneva in July this year, and I quote: *“Most of us underestimate the possibility that the intelligence of AIs will continue to grow. It sounds like science fiction, but it’s a real possibility, and it could change the world in ways that we don’t understand yet, and it could change the power dynamics of our planet in ways that require our attention.”*

This concern is not only about what AI may become in the future, but also about how rapidly it is already being integrated into critical and strategic systems, including the military area, today.

So, in my remarks today, I will try to set out briefly

- Why we need to be attentive to extreme AI risks
- Who needs to be involved
- What needs to be done
- And, not least, how this could happen

The “why” is directly linked to the consequences. Extreme risks are not just words on paper. They are risks that can cause harm on a massive scale, including through the amplification of existing threats such as nuclear, chemical and biological weapons. If a terrorist is able to use AI tools to create bioweapons or to hack into critical systems, what will be the immediate, real-world consequences throughout society and the economy? If one country’s military possesses AI capabilities that allow it to overwhelm an opponent’s defences, who can guarantee that they will resist the temptation to use them? And as AI is increasingly integrated into military decision-support systems, targeting and autonomous functions, how do we ensure that human judgement and responsibility are not reduced to a final “green light” in a process moving too fast for meaningful scrutiny? If an AI system escapes the guardrails that its developers profess to have put in place, who will vouch for what happens next?

As a diplomat by training, I am not given to alarmist statements. But there are enough signals today to justify serious attention and precaution to the possibility that these risks, or others, will emerge. When one of the world's leading authorities on AI says that this is not science fiction, we must listen. When leading companies building frontier AI come together to warn that there is a "limited window" of a few months to strengthen cyber defences and protect against potentially devastating AI-enabled cyberattacks, we must listen. And when AI capabilities are being integrated into military and strategic systems at increasing speed, we should not wait for a crisis before asking ourselves what safeguards are needed. I believe that we should interpret these signals as a clear call for governments and law-makers to act.

The "who needs to be involved" is rather clear. This is a matter for the national security establishment. For the people whose job it is to keep society safe, to identify and manage threats before they turn into actions. But it cannot be left to them alone. To be sure, many actors need to step up at many different levels, including inside parliaments, which have a central role in oversight, legislation and accountability.

Parliaments vary massively in each country as to how they organize their business. All parliaments, nevertheless, can question their government about its approach to extreme AI risks. We are observing that parliaments are starting to establish dedicated committees on AI, or to mainstream AI across the work of different portfolios. We can help to share the experience between parliaments. Parliaments also do not need to start from scratch: existing international legal and normative frameworks can provide important foundations for addressing new AI-related risks, including in the military sphere. And the IPU has already started addressing these issues through various resolutions and the organization of related events. But it is only when the implications for national and international security are clearly understood, that the issue of extreme AI risk will get the necessary political attention and action.

The "what needs to be done" is also rather clear, for better or for worse. The first point is that no single country can manage the risks alone. International cooperation is the only way forward. In our interconnected, digital world, extreme risks could propagate across borders at a speed that will make disasters like the Covid pandemic seem like they belong to another era altogether. This is particularly true where AI amplifies transnational risks such as cyberattacks, biological and chemical threats, or military escalation.

The second point – and here again, I set aside some of my training as a diplomat – is that the United States of America and China will need to agree that it is in their common interest to prevent the development or deployment of AI capabilities whose risks are clearly extreme and unacceptable. I realize that this may sound improbable or impossible to many of you. But I suspect that many people also believe that it is true.

Today, the reality is that the most massive investments in pushing the frontier of AI capabilities are in these two countries. Both see the prize that may be available if they are able to achieve leadership in AI, including military uses of AI. Both see equally clearly the risks if the other country were able to assert leadership, or dominance. So we are locked into a geopolitical struggle, doubled by an intense competition between private companies. Everyone has an incentive to accelerate, including in areas where AI capabilities can have direct implications for military advantage, strategic stability and international security. And the speed of progress on AI capabilities brings us closer every day to what Professor Bengio reminds us is no longer science fiction.

But this growing awareness of extreme risk also creates an opportunity, including for the USA and China, to pull back where the risks are greatest. To slow the pace of frontier AI development where necessary. To establish red lines that AI must not be allowed to cross, including in areas where its use could undermine strategic stability, meaningful human control or the protection of civilians. To create obligations for private companies to demonstrate that their products are safe, just as is done in so many other domains. Airplanes must be safe. Nuclear reactors must be safe. Pharmaceuticals must be safe. AI must be safe, too.

I like to imagine that many of you may agree with this sentiment, while also thinking that a shared commitment to restraint in the most dangerous areas sounds like science fiction. It is our job to make sure that it becomes a reality.

Which brings me to my final point, the “how this could happen”. Many people in this room have deep expertise in these questions, but from a political and diplomatic perspective, I would like to offer three pointers: the role of the middle powers, the United Nations, and parliamentary diplomacy.

It escapes no-one’s attention that the great powers that have previously underpinned global norms have been taking steps to undermine the very norms that they created. These unwelcome developments place a new responsibility on middle powers to come together to lead, as we have heard from the Canadian Prime Minister and very recently former UK Prime Minister Gordon Brown, among other voices. The combined efforts of these countries can help to build a shared understanding that no-one gains from allowing extreme AI risks to develop unchecked, laying the groundwork for others to join.

The United Nations was established after the Second World War to prevent such widespread destruction from ever happening again. Today, that system is under huge financial pressures and, more importantly, under growing political pressure from Member States themselves. If we believe in multilateralism, we have to work through the United Nations as the legitimate forum for establishing global norms. A new Secretary-General will be elected shortly. We must ensure that she (or he) has the mandate to bring all parties to the table to work towards AI that serves humanity and protects international peace and security.

Lastly, I respectfully submit that parliamentary diplomacy is a potential channel for this necessary dialogue. Exchanges between parliamentarians can sometimes help to unlock doors at moments of profound political tension. We saw this during the Cold War, when the inter-parliamentary conference on peace and security in Europe brought together parliamentarians from East and West, leading to the creation of the OSCE. More recently, we observe the constructive engagement at the IPU between Armenia and Azerbaijan.

The IPU was founded in 1889 with the simple idea that dialogue between parliamentarians can provide a space for building trust and mutual understanding, addressing political differences and finding solutions to complex problems. I believe that this idea remains just as relevant today. I want to believe that inter-parliamentary dialogue on extreme AI risks is possible, a dialogue that involves all stakeholders, including the USA and China, the middle powers and everyone who is ready to step forward and work to ensure the safe and responsible development of AI.

And this is perhaps where our essential convergence can begin: by recognizing that, whatever our different perspectives on the future of AI, we share an interest in preventing its most dangerous applications from threatening people, international peace and security, and our common humanity.

Thank you.