

Ahli-CEO Ingatkan Potensi Kepunahan, Samakan AI dengan Perang Nuklir

Category: Teknologi

written by Maulya | 31/05/2023



[Orinews.id](https://orinews.id)| **Jakarta** – Para ahli dan pimpinan perusahaan (CEO) melontarkan peringatan kedua terkait bahaya kepunahan gara-gara kecerdasan buatan (AI).

Hal itu terungkap dalam pernyataan terbuka ‘Statement on AI Risk’ yang digagas oleh organisasi nirlaba Center for AI Safety yang berbasis di San Francisco, AS.

Pendatangannya mencapai lebih dari 100 tokoh yang merupakan ilmuwan, termasuk Geoffrey Hinton dan Yoshua Bengio yang memenangkan Turing Award 2018 atas karya mereka di bidang kecerdasan buatan.

Selain itu, ada pula para pimpinan perusahaan teknologi, termasuk CEO Google DeepMind Demis Hassabis dan CEO OpenAI yang merupakan pemilik chatbot ChatGPT, Sam Altman.

“Mitigasi risiko kepunahan dari AI harus menjadi prioritas global bersama dengan risiko skala masyarakat lainnya, seperti pandemi dan perang nuklir,” demikian tertulis pada pernyataan terbuka itu.

Pernyataan tersebut merupakan intervensi terbaru dalam perdebatan yang rumit dan kontroversial mengenai keamanan AI.

Awal tahun ini, sebuah surat terbuka sejenis yang ditandatangani oleh beberapa tokoh yang kurang lebih sama dengan yang mendukung peringatan terbaru ini menyerukan “jeda” selama enam bulan dalam pengembangan AI.

Namun, seruan tersebut dikritik dalam berbagai tingkatan. Beberapa ahli menganggap surat tersebut melebih-lebihkan risiko yang ditimbulkan oleh AI, sementara yang lain setuju dengan risiko tersebut, tetapi tidak dengan solusi yang disarankan dalam surat tersebut.

Direktur eksekutif Center for AI Safety Dan Hendrycks mengatakan pernyataan terbaru yang dilontarkan secara singkat ini dimaksudkan untuk menghindari ketidaksepakatan tersebut.

“Kami tidak ingin mendorong menu yang sangat besar dengan 30 intervensi potensial,” kata Hendrycks, dikutip dari The Verge.

“Ketika hal itu terjadi, maka akan melemahkan pesan yang ingin disampaikan,” tambahnya.

Hendrycks menggambarkan pesan tersebut sebagai “coming-out” bagi para tokoh di industri yang khawatir akan risiko AI.

“Ada kesalahpahaman yang sangat umum, bahkan di komunitas AI, bahwa hanya ada segelintir orang yang khawatir pada hal tersebut. Namun, pada kenyataannya, banyak orang yang secara pribadi mengungkapkan kekhawatirannya tentang hal-hal ini,” ujar Hendrycks.

Kontroversi soal AI sendiri sudah tidak asing, tetapi pembicaraan soal detailnya sering kali tak berkesudahan. Namun

intinya, ada hipotetis di mana sistem AI dengan cepat meningkatkan kemampuannya, dan tidak lagi berfungsi dengan aman.

Banyak ahli merujuk pada peningkatan yang cepat dalam sistem seperti model bahasa yang besar sebagai bukti proyeksi peningkatan kecerdasan di masa depan.

Mereka mengatakan setelah sistem AI mencapai tingkat kecanggihan tertentu, mungkin akan menjadi tidak mungkin untuk mengendalikan tindakan mereka.

Namun, beberapa pihak meragukan prediksi ini. Mereka menunjukkan ketidakmampuan sistem AI untuk menangani tugas-tugas yang relatif biasa seperti, misalnya, mengemudikan mobil.

Meski sudah bertahun-tahun dan miliaran investasi di bidang penelitian ini, mobil yang sepenuhnya bisa menyetir sendiri masih jauh dari kenyataan.

Menurut mereka, jika AI tidak dapat menangani tantangan yang satu ini, maka peluang yang dimiliki teknologi ini untuk menyamai setiap pencapaian manusia di tahun-tahun mendatang cukup kecil.

Meski demikian, baik pendukung risiko AI maupun yang skeptis setuju jika sistem AI tidak ditingkatkan maka akan menghadirkan sejumlah ancaman di masa sekarang.

Contohnya, penggunaan untuk pengawasan massal di dunia maya, hingga kemudahan penciptaan misinformasi dan disinformasi.

|Sumber: CNNIndonesia