



INTERACTIONS BETWEEN THE AI CONTROL PROBLEM AND THE GOVERNANCE PROBLEM

Professor Nick Bostrom

University of Oxford

Director, Strategic Artificial Intelligence Research Center

Director, Future of Humanity Institute

Q:

“If human level general AI is developed, then what are likely outcomes?”

Google DeepMind Challenge Match 4국 ILVEI 바둑TV

이세돌 9단 VS 알파고 2

01:59:42

01:59:42

SAM IK



Lee Sedol v. AlphaGO

Oct 2015

“Based on its level seen in the match (against Fan), I think I will win the game by a near landslide”

Feb 2016

“I have heard that Google DeepMind’s AI is surprisingly strong and getting stronger, but I am confident that I can win at least this time”

Mar 9, 2016

“I was very surprised because I didn’t think I would lose”

Mar 10, 2016

“I’m quite speechless ... I am in shock. I can admit that ... the third game is not going to be easy for me”

Mar 12, 2016

“I kind of felt powerless.”

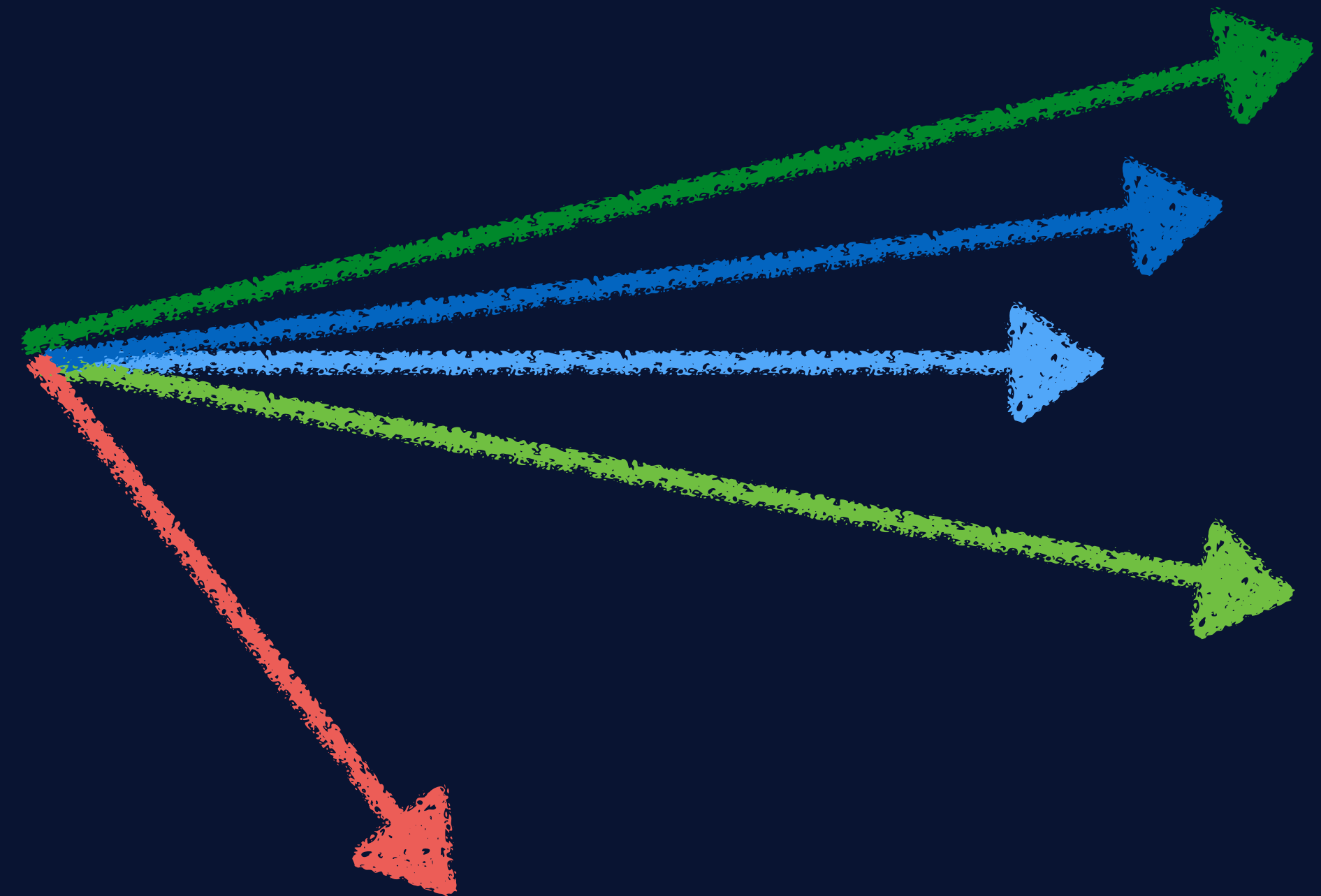


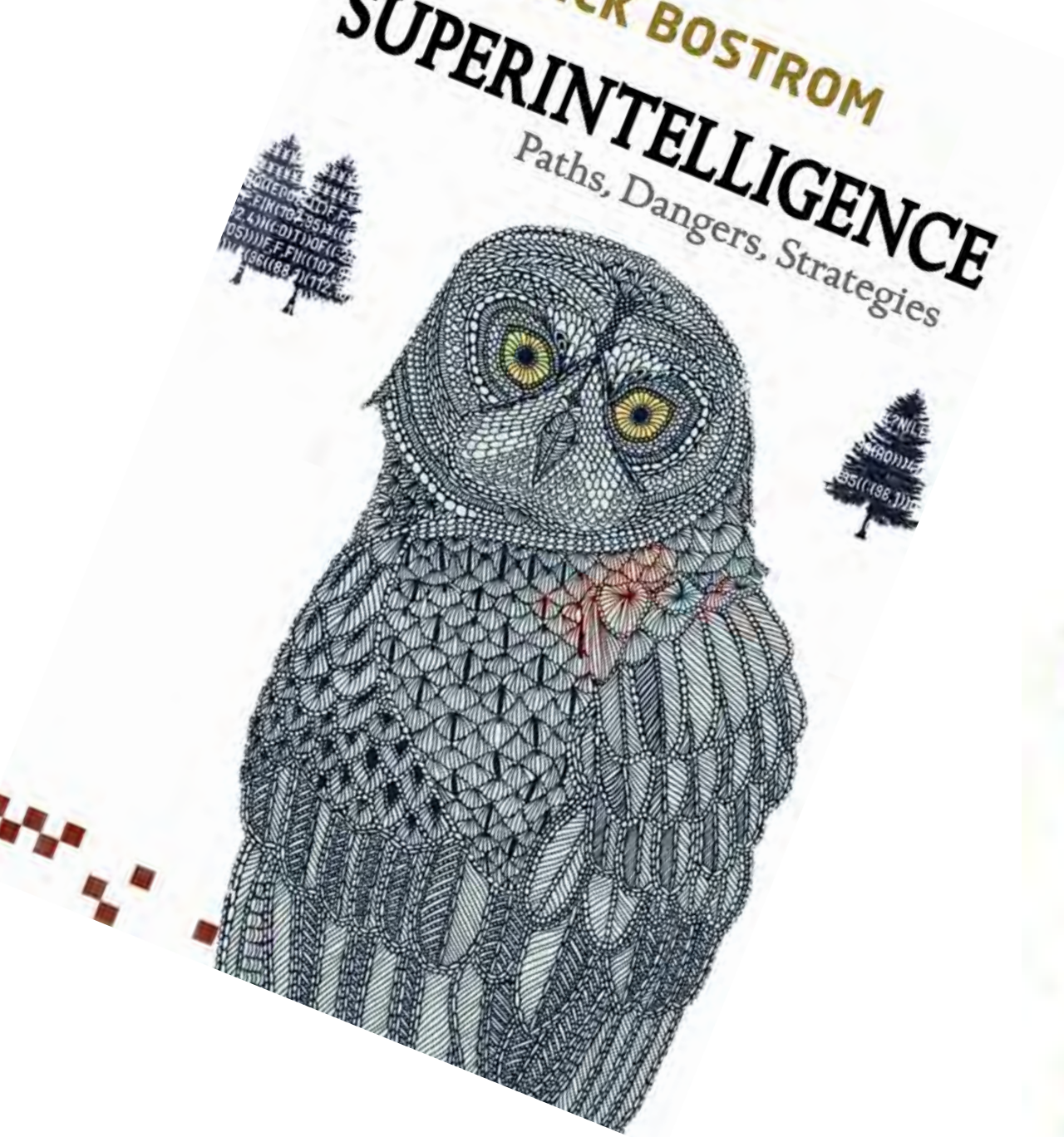
Q:

“If human level general AI is developed, then what are likely outcomes?”

A:

Superintelligence





0 of 1 people found the following review helpful



One Star, September 10, 2016

By [Amazon Customer](#)

Verified Purchase ([What's this?](#))

This review is from: **Superintelligence: Paths, Dangers, Strategies (Paperback)**

I read it in 3 days and I'm profoundly depressed.

Help other customers find the most helpful reviews

Was this review helpful to you?

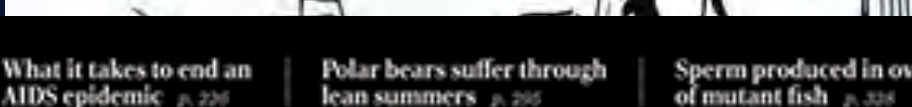
[Report abuse](#) | [Permalink](#)

Q:

“What can we do now to maximize the probability of a positive outcome?”

A:

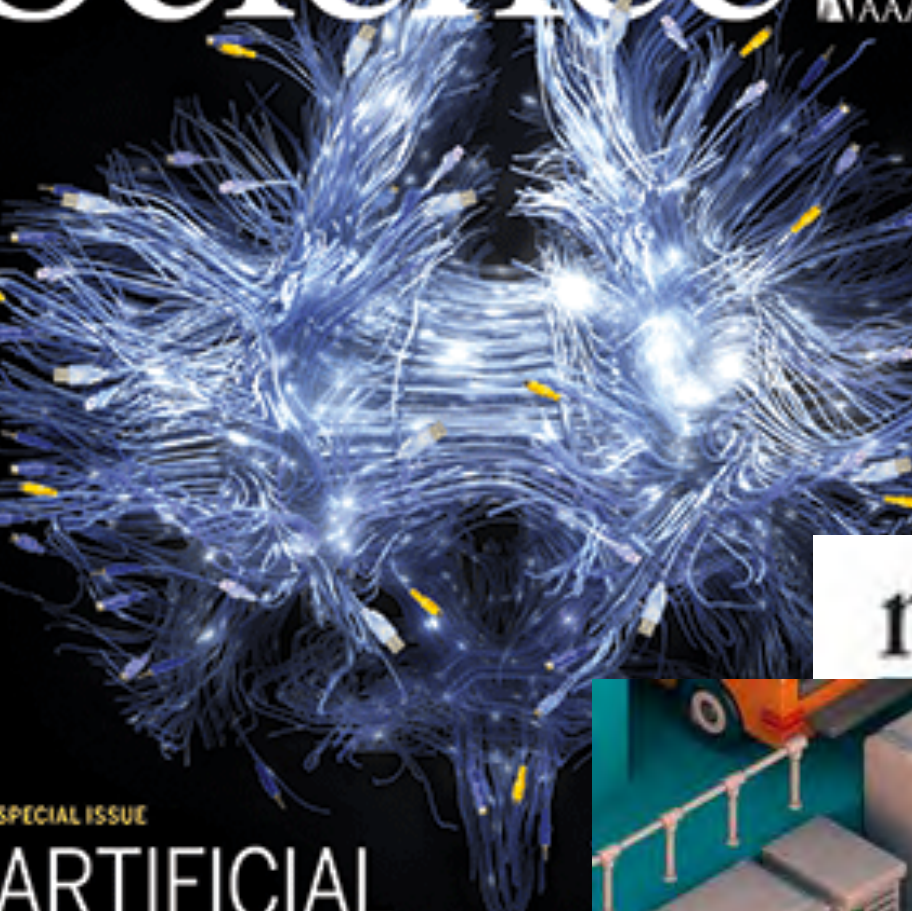
- 
- solve intelligence
 - solve scalable control
 - solve AI governance problem



Science

\$10
17 JULY 2015
sciencemag.org

AAAS



SPECIAL ISSUE
ARTIFICIAL INTELLIGENCE

AREA CHAIRS

Jacob Aberkane (University of Michigan), Ryan Adams (Harvard U), University of California Berkeley), Samy Bengio (Google), Yoshua B. Coss-Granchi (Università degli Studi di Milano), Kamada Chutaba (Tsukuba), Marco Cuturi (Kyoto University), Trevor Darrell (Stanford), Felipe Fernandez (Universitat de València), Theodor Gneiting (University of Cologne), Johannes Heule (University of Texas at Austin), Mayur Gupta (Google), Kristen Grauman (University of Texas at Austin), Mounir Hachimi (University of Montreal), Sanyan Kala (Yahoo! Labs), Balazs Kegl (University of Waterloo), Simon Lacoste-Julien (HRI/H), Hugo Larochette (Twitter/Toronto University), Phil Long (Swinburne Technology), Sher Manassu (University College London), Oleks Prokop (McGill University), Ma (University of Basel), Daniel Roy (University of Toronto), S.W.K. Wei, Matthias Seeger (Amazon Research), Fei Sha (University of Southern California), Michaela Institute of Technology), Karl Schindler (Cornell University of Alberta), Ralf Stein (Max Planck Institute Tübingen), Raquel Urtasun (University of British Columbia), Corinna Voelet (University of Zurich), Eric Xing (Carnegie Mellon University), Yuesong Yu (Tsinghua), Zhen Zhong (Harbin)

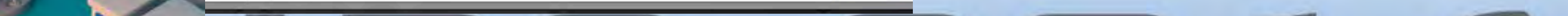
www.icml.cc/2016/



Learning



undage @Miles_Brundage · Jan 1
on Recognition Based on Joint Trajectory
s with Convolutional Neural Networks,"
g et al.: arxiv.org/abs/1612.09401



BARCELONA · SPAIN · DECEMBER 5 - 10, 2015 | <http://nips.cc/>

technical research agendas

- inverse reinforcement learning
- adversarial examples
- models of control failure
- approval-maximizing agents
- imitation agents
- architectural composition
- corrigibility
- foundations of reflective agents
- detecting context change
- interpretability and explanation
- control diversification



the AI governance problem

Strategic Implications of Openness in AI Development¹

(2016) Nick Bostrom
Future of Humanity Institute
Strategic Artificial Intelligence Research Centre
University of Oxford

www.nickbostrom.com
[Global Policy (2017), in press]

Abstract

This paper attempts a preliminary analysis of the global desirability of different forms of openness in AI development (including openness about source code, science, data, safety techniques, capabilities, and goals). Short-term impacts of increased openness appear mostly socially beneficial in expectation. The strategic implications of medium- and long-term impacts are complex. The evaluation of long-term impacts, in particular, may depend on whether the objective is to benefit the present generation or to promote a time-neutral aggregate of well-being of future generations. Some forms of openness are plausibly positive on both counts (openness about safety measures, openness about goals). Others (openness about source code, science, and possibly capability) could lead to a tightening of the competitive situation around the time of the introduction of advanced AI, increasing the probability that winning the AI race is incompatible with using any safety method that incurs a delay or limits performance. We identify several key factors that must be taken into account by any well-founded opinion on the matter.

Background

The goal of this paper is to conduct a preliminary analysis of the long-term strategic implications of openness in AI development. What effects would increased openness in AI development have, on the margin, on the long-term impacts of AI? Is the expected value for society of these effects positive or negative? Since it is typically impossible to provide definitive answers to this type of question, our ambition here is more modest: to introduce some relevant considerations and develop some thoughts on their weight and plausibility.

¹ For helpful comments and discussion, I'm grateful to Stuart Armstrong, Owen Cotton-Barratt, Rob Bensinger, Miles Brundage, Paul Christiano, Allan Dafoe, Eric Drexler, Owain Evans, Oliver Habryka, Denis Hassabis, Shane Legg, Javier Lezama, Luke Muehlhauser, Toby Ord, Guy Riviere, Steve Raynor, Anders Sandberg, Andrew Simpson, and Mustafa Sulayman. I'm especially grateful to Carrick Flynn and Carl Shulman for help with several parts of the manuscript.

Policy Desiderata in the Development of Machine Superintelligence¹

(2016) version 3.4

Nick Bostrom
Allan Dafoe
Carrick Flynn

[working paper]
www.nickbostrom.com

ABSTRACT

Machine superintelligence could plausibly be developed in the coming decades or century. The prospect of this transformative development presents a host of political challenges and opportunities. This paper seeks to initiate discussion of these by identifying a set of distinctive features of the transition to a machine intelligence era. From these distinctive features, we derive a correlative set of policy desiderata—considerations that should be given extra weight in long-term AI policy compared to other policy contexts. We argue that these desiderata are relevant for a wide range of actors (including states, AI technology firms, investors, and NGOs) with an interest in AI policy. However, developing concrete policy options that satisfy these desiderata will require additional work.

¹ For comment and discussion, we're grateful to Stuart Armstrong, Michael Barnett, Seth Baum, Andrew Snyder-Beattie, Dominic Becker, Nick Beckstead, Devi Borg, Miles Brundage, Paul Christiano, Jack Clark, Rebecca Crocker, Richard Danzig, Daniel Dewey, Eric Drexler, Sebastian Farquhar, Sophie Fischer, Ben Garfinkel, Tom Grant, John Halstead, Robin Hanson, Verity Harding, Sean Legassick, Wendy Lin, Jelena Luketina, Luke Muehlhauser, Toby Ord, Mahendra Prasad, Anders Sandberg, Carl Shulman, Nate Soares, Mojmir Stehlik, and Alex Tymchenko.

AI & Soc (2016) 31:201–208
DOI 10.1007/s00146-015-0590-y



OPEN FORUM

Racing to the precipice: a model of artificial intelligence development

Stuart Armstrong¹ · Nick Bostrom¹ · Carl Shulman¹

Received: 29 April 2015 / Accepted: 6 July 2015 / Published online: 1 August 2015
© Springer-Verlag London 2015

Abstract This paper presents a simple model of an AI (artificial intelligence) arms race, where several development teams race to build the first AI. Under the assumption that the first AI will be very powerful and transformative, each team is incentivised to finish first—by skimping on safety precautions if need be. This paper presents the Nash equilibrium of this process, where each team takes the correct amount of safety precautions in the arms race. Having extra development teams and extra enmity between teams can increase the danger of an AI disaster, especially if risk-taking is more important than skill in developing the AI. Surprisingly, information also increases the risks: the more teams know about each others' capabilities (and about their own), the more the danger increases. Should these results persist in more realistic models and analysis, it points the way to methods of increasing the chance of the safe development of AI.

Keywords AI · Artificial intelligence · Risk · Arms race · Coordination problem · Model

1 Introduction

This paper presents a simplified model for analysing technology races. The model was designed initially for races to construct artificial intelligences (AIs). But it can be applied to other similar races or competitions, especially

technological races where there is a large advantage to reaching the goal first.

There are arguments that the first true AIs are likely to be extremely powerful machines (Good 1965; Chalmers 2010), but that they *could* end up being dangerous (Omohundro 2008; Yudkowsky 2008; Bostrom 2014) if not carefully controlled (Armstrong et al. 2012).

This paper is too short to discuss these AI disaster scenarios—there are many different scenarios and factors to consider (Bostrom 2014), including subtle issues such as the orthogonality of value and intelligence (Bostrom 2012; Armstrong 2013), and the uncertainties in prediction are extremely high (Armstrong and Soata 2012; Armstrong et al. 2014)¹. The simplified AI risk thesis is that AIs could become extremely powerful and then act to reshape the world according to their preferences, and these preferences might not be human friendly² (thus various research projects such as "friendly AI" (Muehlhauser and Salamon 2012) trying to ensure the creation of an AI with human-compatible values). Such powerful entities could be extremely difficult to control or turn-off at that point (as it would almost certainly be motivated to resist such control (Omohundro 2008)). However, such dangers are certainly not within the capabilities of current automation (see the expert poll at Sandberg and Bostrom (2011)); thus, we will avoid discussion of recent developments in the field.

So in this paper, we will consider a simplified and general scenario. We will assume that there is a definite probability of an AI-related disaster, given the creation of AI. We will further assume that the probability of such

¹ Though high uncertainties do not imply safety.

² Some examples of bad values could be "prevent human suffering" via killing people off and "make people happy" via wireheading (Yudkowsky 2008; Bostrom 2014).

Openness



- safety measures ✓
- values ✓
- (capability)
- source code, platforms
- science
- training data, environments, benchmarks

} ?

Observation

Openness reduces the gap between the leading developer and the nearest follower.

- a couple of years in a low openness scenario?
- a few months in a high openness scenario?
- zero in the limiting case of maximal openness

This could help reduce the risk that a small group monopolizes all the benefits.

Suppose that...



- safety requires some significant extra work after AI is completed

OR

→ doom

- safe operation initially incurs a significant performance penalty

OR

→ doom

- the Vulnerable World Hypothesis is true in the post-AI-transition world

Vulnerable world hypothesis

There is some level of technology at which offense strongly dominates defense, in the sense that any small group of reasonably competent people with access to the technology would be able to take some action that would lead to the destruction of the world (independently of what other people did after the action was taken).

biotechnology?

nanotechnology?

doomsday device?

Suppose that...



- safety requires some significant extra work after AI is completed

OR

→ doom

- safe operation initially incurs a significant performance penalty

OR

→ doom

- the Vulnerable World Hypothesis is true in the post-AI-transition world

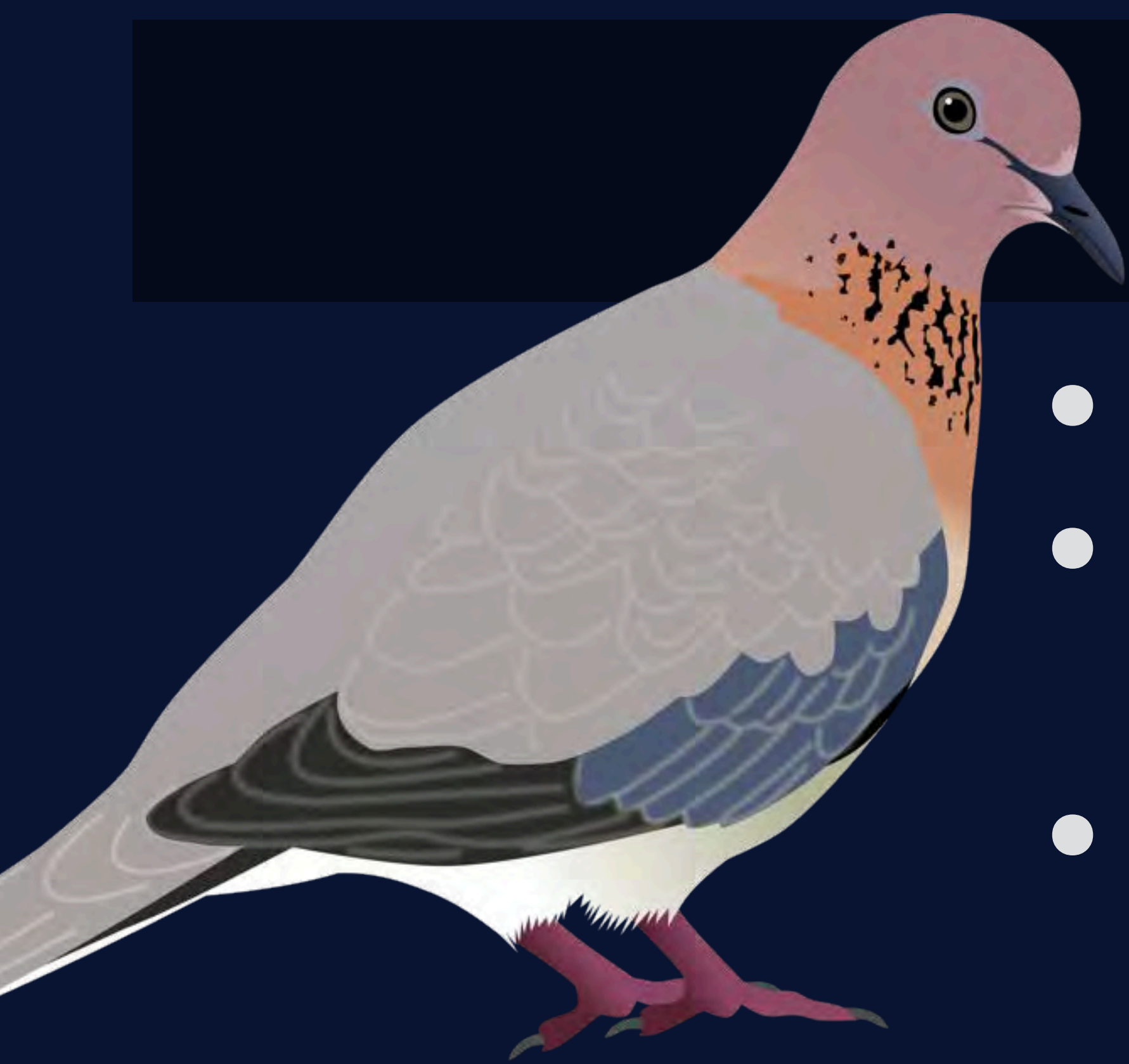
Openness



- safety measures ✓
- values ✓
- (capability)
- source code, platforms
- science
- training data, environments, benchmarks



What to do?



- openness for now
- desired property: **conditional stabilization...**
- lay the foundations for a collaborative approach later:
 - coordinate (or ideally pool) research among trusted leading groups
 - create ability **not** to share science and algorithms until it is safe to do so
 - credibly commit to sharing benefits and influence

STRATEGIC
ARTIFICIAL
INTELLIGENCE
RESEARCH
CENTRE



Future
of Humanity
Institute



UNIVERSITY OF OXFORD