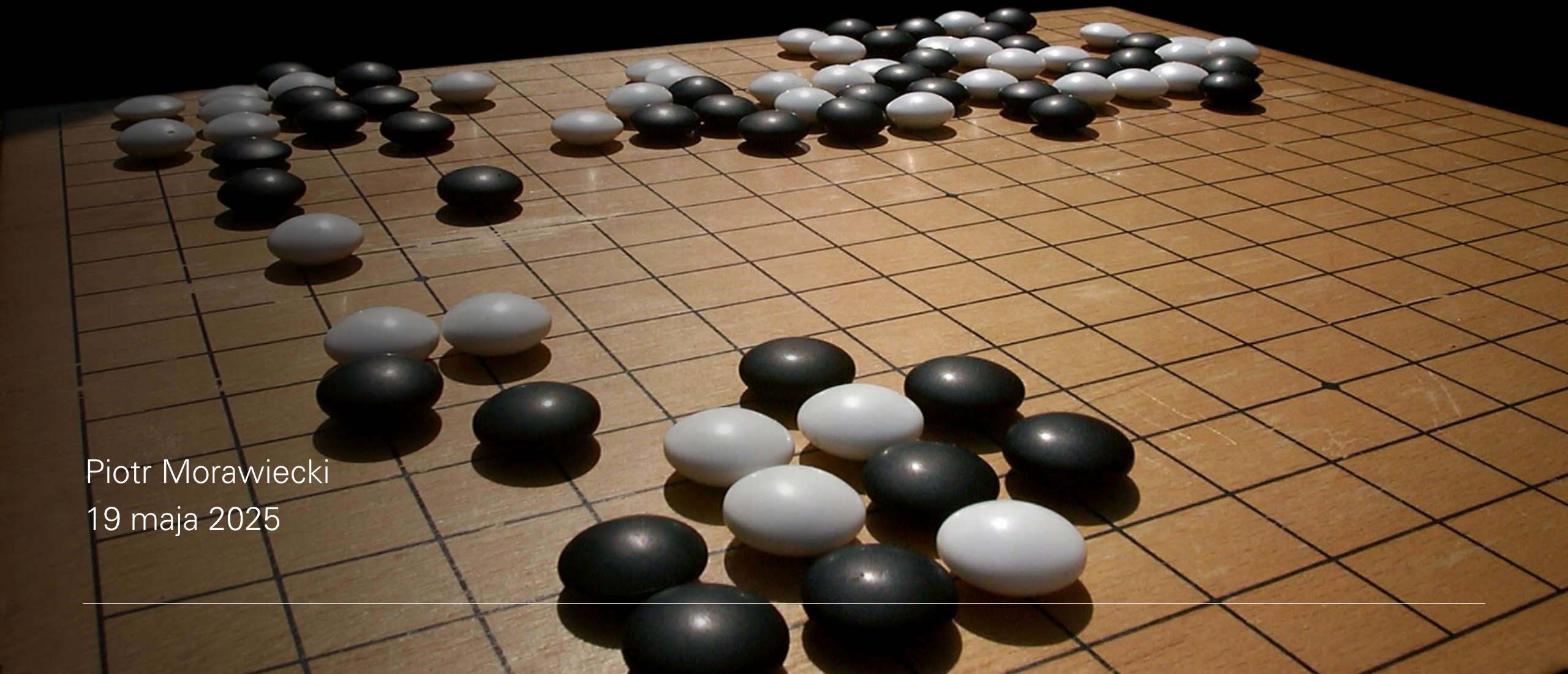
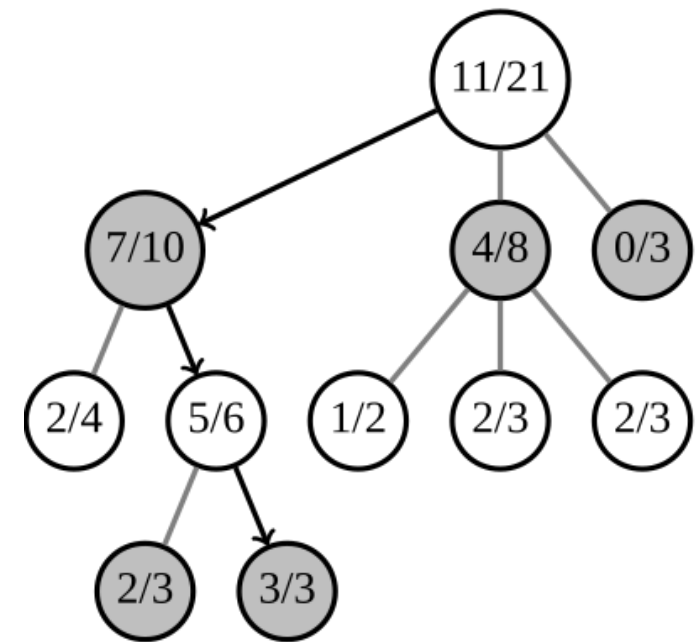

Wykład 10. O rywalizacji człowieka z komputerem

Piotr Morawiecki
19 maja 2025



Podsumowanie poprzedniego wykładu

- Poznaliśmy algorytm *Monte Carlo Tree Search* (MCTS), który pozwala znajdować strategie w grach, w których przeszukanie całego drzewa gry jest niemożliwe.
- Polega na iteracyjnym przeszukiwaniu najbardziej obiecujących gałęzi drzewa gry, a następnie na drodze symulacji ocenianiu szans na wygraną przy wyborze danej strategii.
- Algorytmu MCTS można wzmocnić wykorzystując metody heurystyczne, np. ewaluacją pośrednich stanów gry.



Dwa słynne pojedynki człowieka z komputerem

DEEP-BLUE VS GARY KASPAROV (1997)



ALPHA GO VS LEE SEDOL (2016)



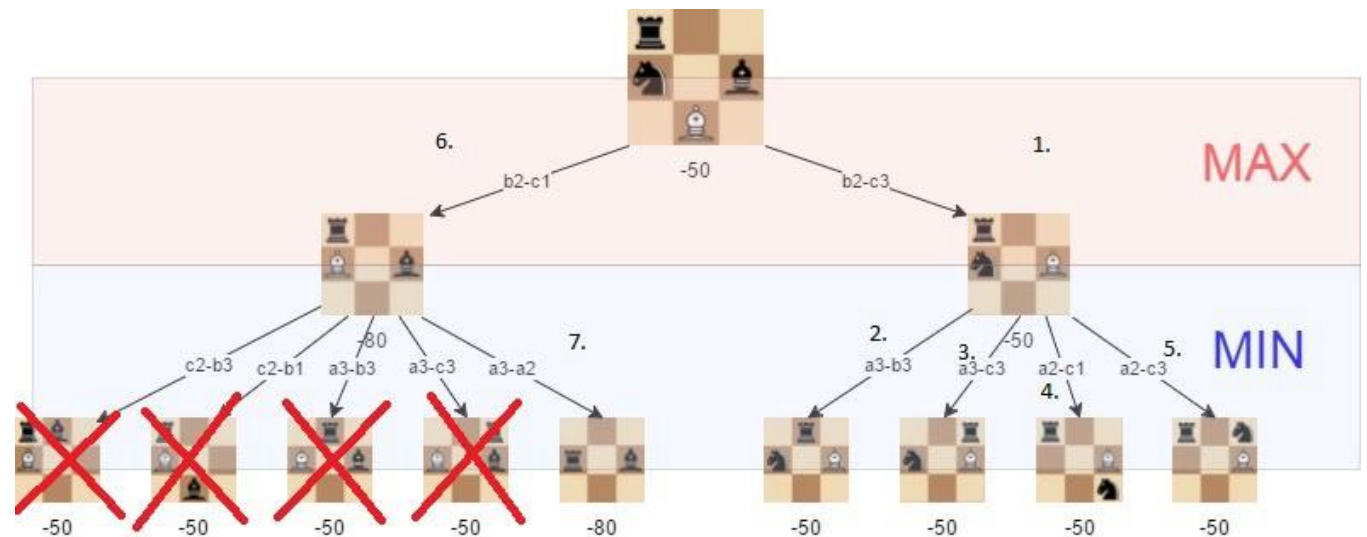
Algorytm DeepBlue (1/2)

- DeepBlue został zaprogramowany w oparciu o wiedzę ekspercką.
- Wykorzystywał on bazę danych zawierającą zapis 700,000 rozgrywek arcymistrzów szachowych.
- W oparciu o nią skonstruował zbiór danych (tzw. *Extended Book*), który zawierał miliony możliwych stanów gry z pierwszych 30 ruchów, pozwalając algorytmowi wybierać pierwsze ruchy.
- Ponadto zawierał obszerną bazę zakończeń gry (tzw. *Endgame Database*), zawierającą wynik gry po tym jak na planszy zostanie 6 lub mniej figur szachowych.



Algorytm DeepBlue (2/2)

- Podczas rozgrywki wykorzystywał algorytm minimax dla kilku następnych poziomów drzewa gry. Stany gry były oceniane z wykorzystaniem heurystycznych wzorów uwzględniających:
 - wartości figur każdego gracza,
 - położenie figur (bezpieczeństwo króla, ustawienie pionków, aktywność figur, itd.), oraz
 - informacje z *endgame database*.
- Superkomputer analizował ok. 200 milionów stanów gry na sekundę.



DeepBlue vs Gary Kasparov

- DeepBlue zmierzył się w 1996 roku z mistrzem świata w szachach, Garrym Kasparovem. DeepBlue wygrał pierwszą rozgrywkę, ale później zremisował dwóch i przegrał w trzech.
- Rewanż odbył się w 1997 roku, podczas którego Deep Blue wygrał dwie rozgrywki, Garry Kasparov jedną, a pozostałe trzy zakończyły się remisem.



Czas na nowe wyzwanie...

„The ancient oriental game of Go has long been considered a grand challenge for artificial intelligence. For decades, computer Go has defied the classical methods in game tree search that worked so successfully for chess and checkers.”

S. Gelly, et al. *The grand challenge of computer Go: Monte Carlo tree search and extensions*. (2012)

Zasady Go

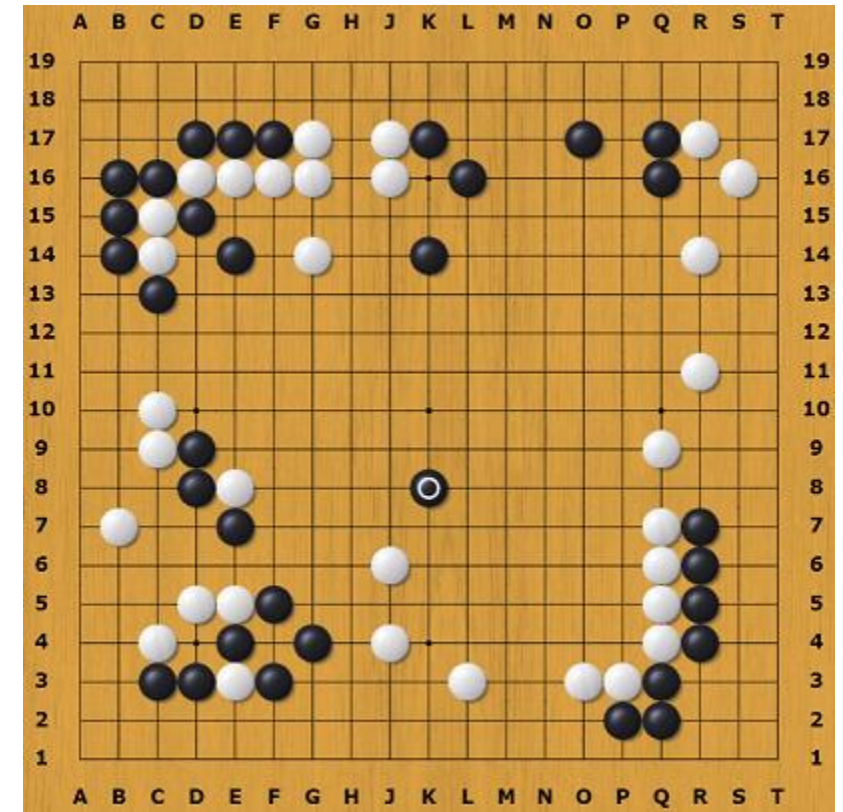


Na czym polega trudność Go?

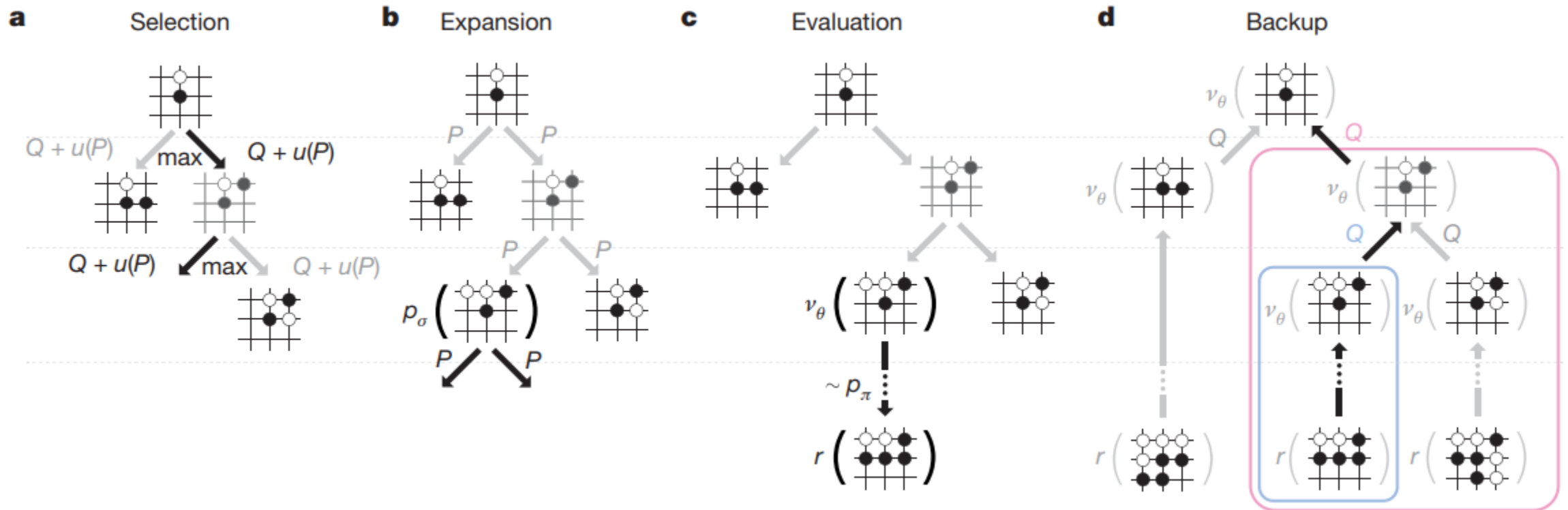
1. Ogromna liczba możliwych ruchów

	Szachy	Go
Liczba kombinacji dwóch pierwszych ruchów	$20^2 = 400$	$361 \cdot 360 \approx 130,000$
Typowa liczba tur w trakcie jednej partii	40 – 60	≈ 200
Szacowana liczba możliwych stanów gry	10^{47}	10^{170}

2. Trudność dokonaniem ewaluacji stanu gry (każdy ruch może mieć niespodziewane konsekwencje na dużo późniejszym etapie gry)



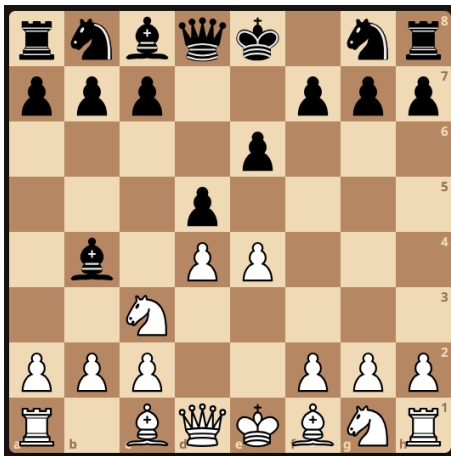
Algorytm AlphaGo: MCTS



Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Czym są sieci neuronowe?

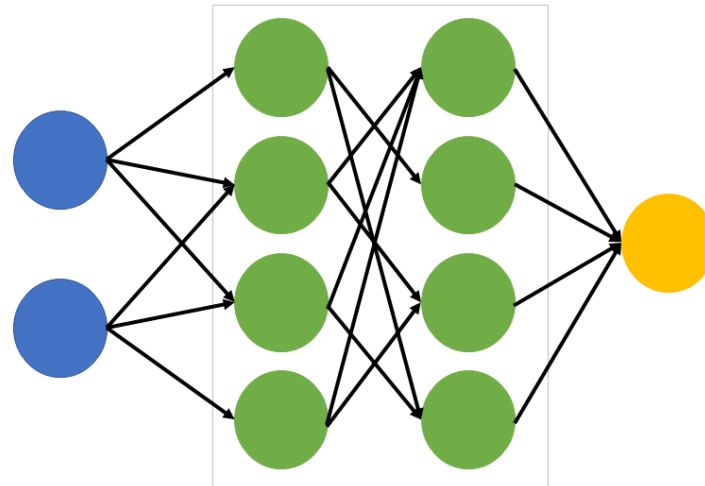
Wejście: stan gry



Input Layer

Hidden Layer

Output Layer



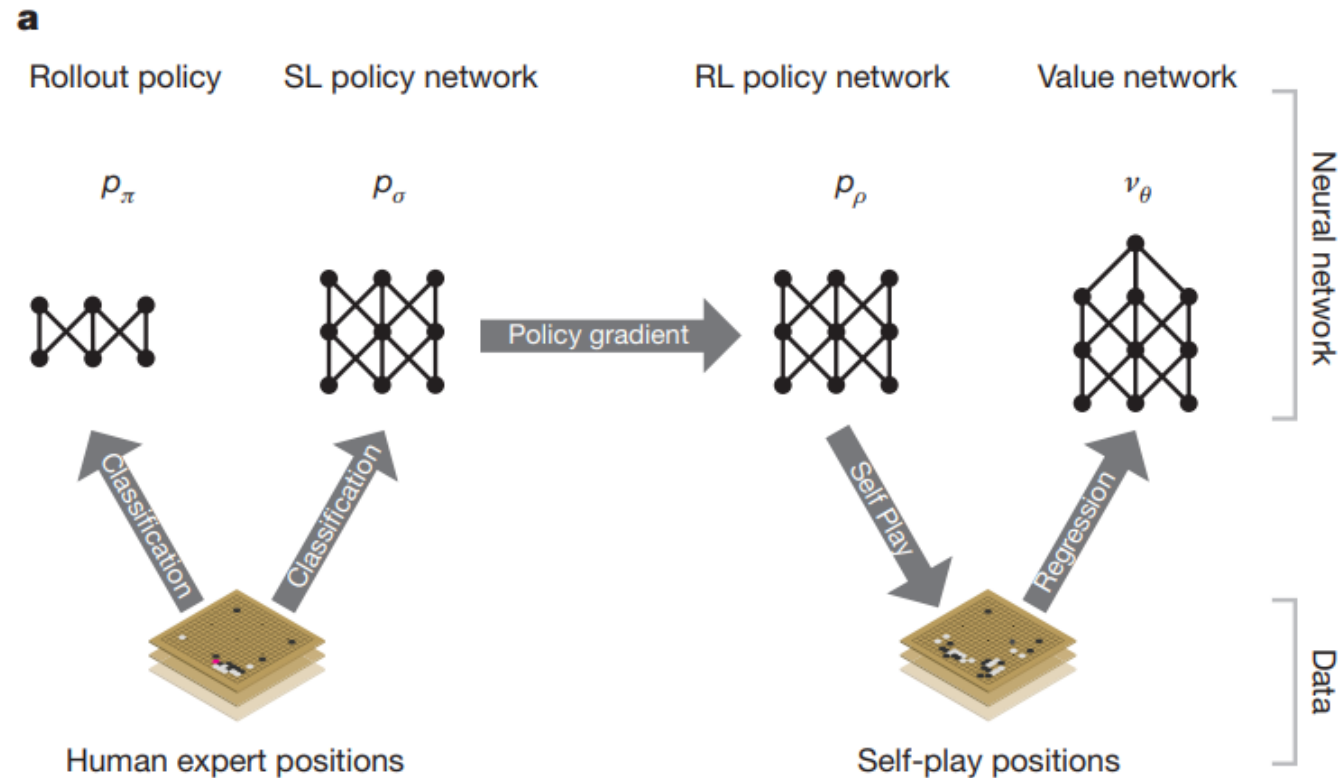
Artificial Neural Networks

Wyjście:

prawdopodobieństwo
wygranej danego gracza

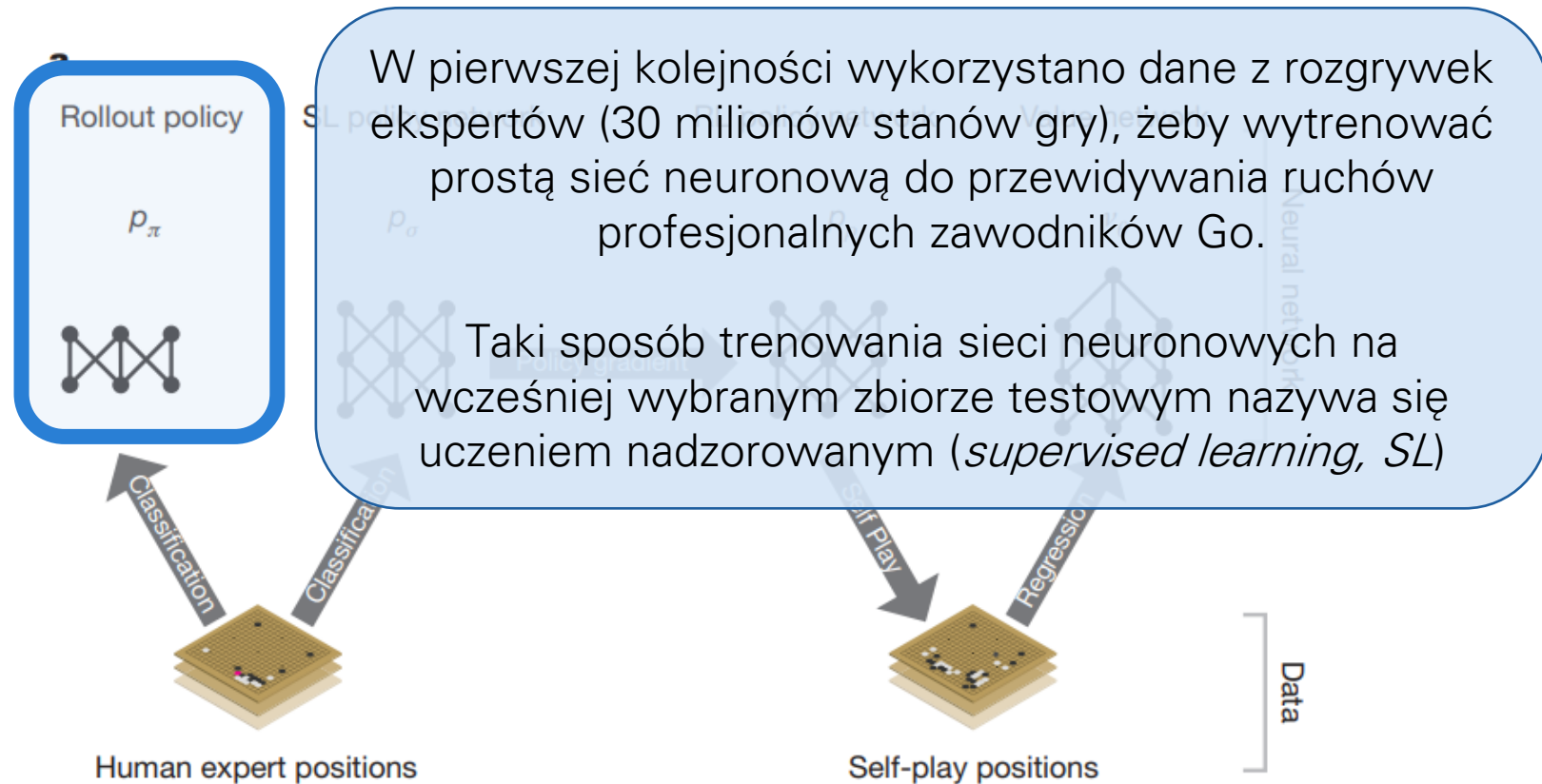
Interaktywna demonstracja: <http://playground.tensorflow.org/>

Algorytm AlphaGo: trenowanie sieci neuronowych



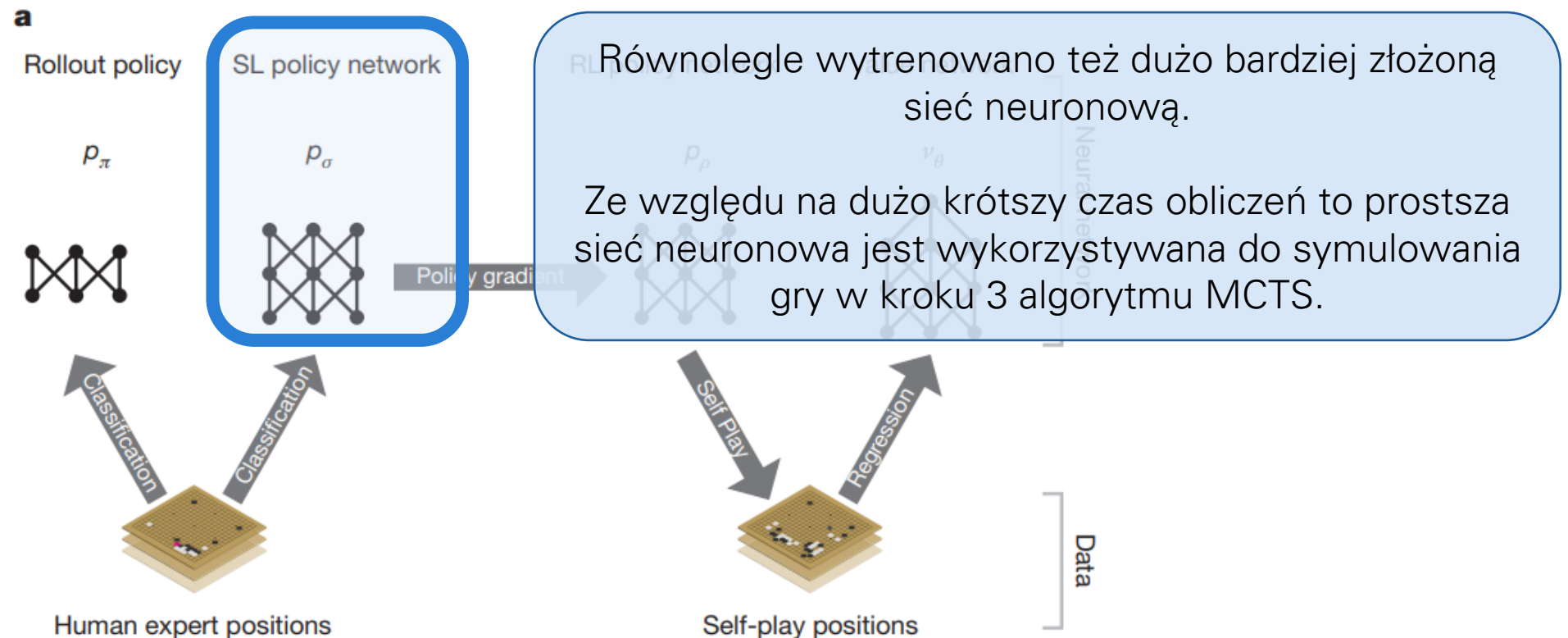
Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: trenowanie sieci neuronowych



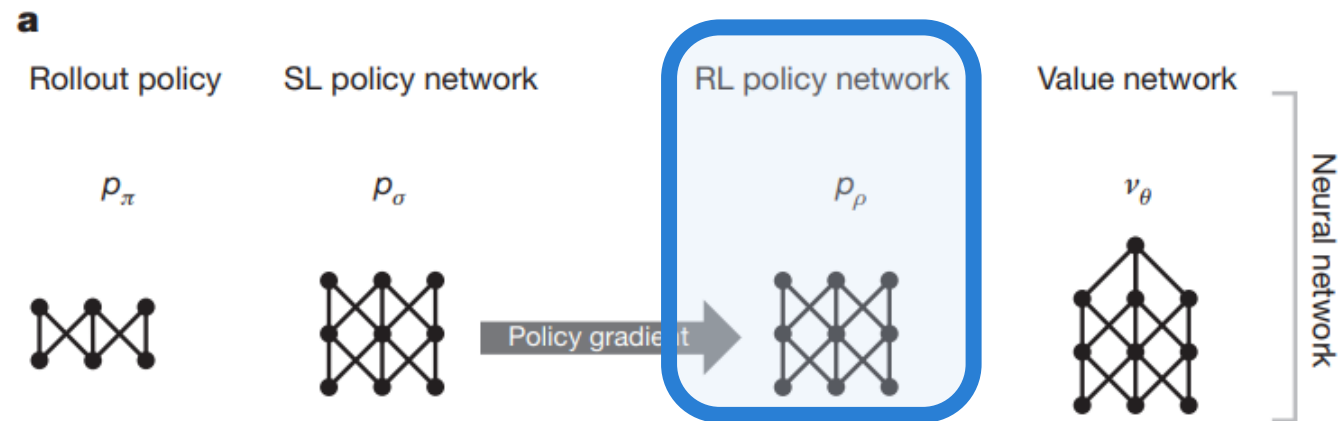
Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: trenowanie sieci neuronowych



Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: trenowanie sieci neuronowych

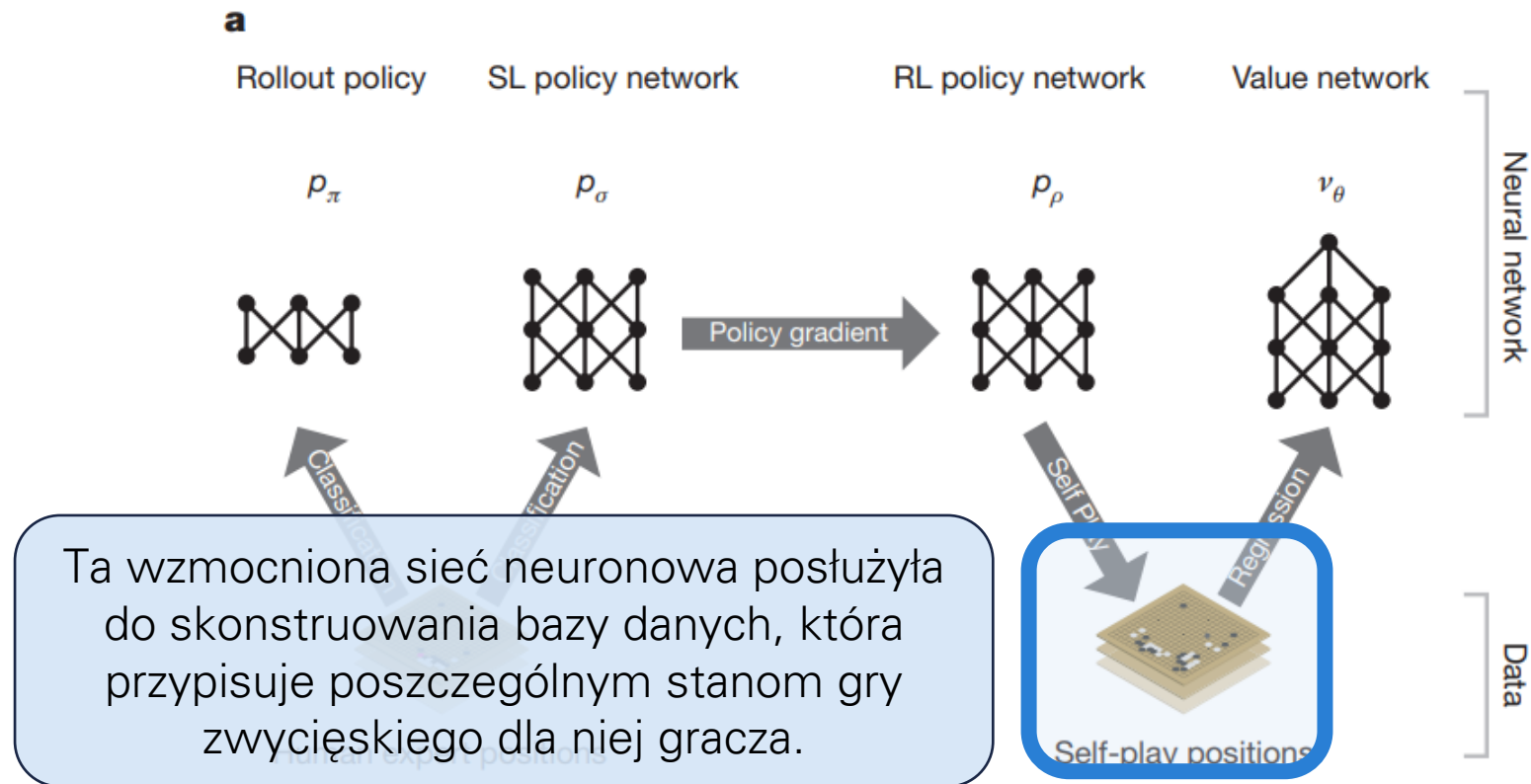


Bardziej złożona sieć została dalej udoskonalona z wykorzystaniem tzw. uczenia wzmocnionego (*reinforcement learning, RL*).

Algorytm rozgrywał partie Go sam ze sobą, za każdym razem modyfikując wagi poszczególnych neuronów, tak żeby zmaksymalizować szansę na wygraną.

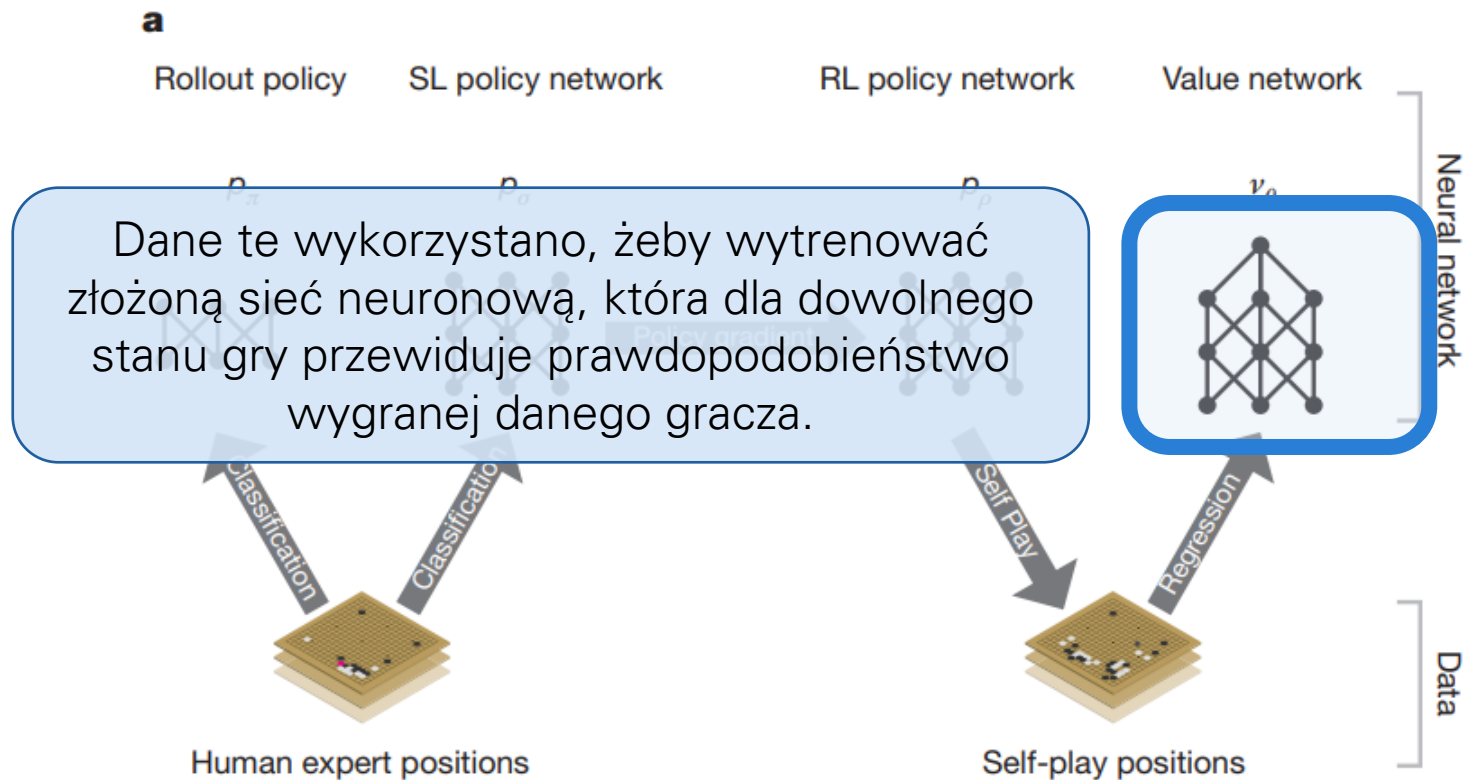
Źródło: D. Silver, A. Huang, C. Madhavan, J. Schaeffer, A. Tang, D. Zeman, C. Bowling, *Starcraft: A Challenge for General Game Playing*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: trenowanie sieci neuronowych



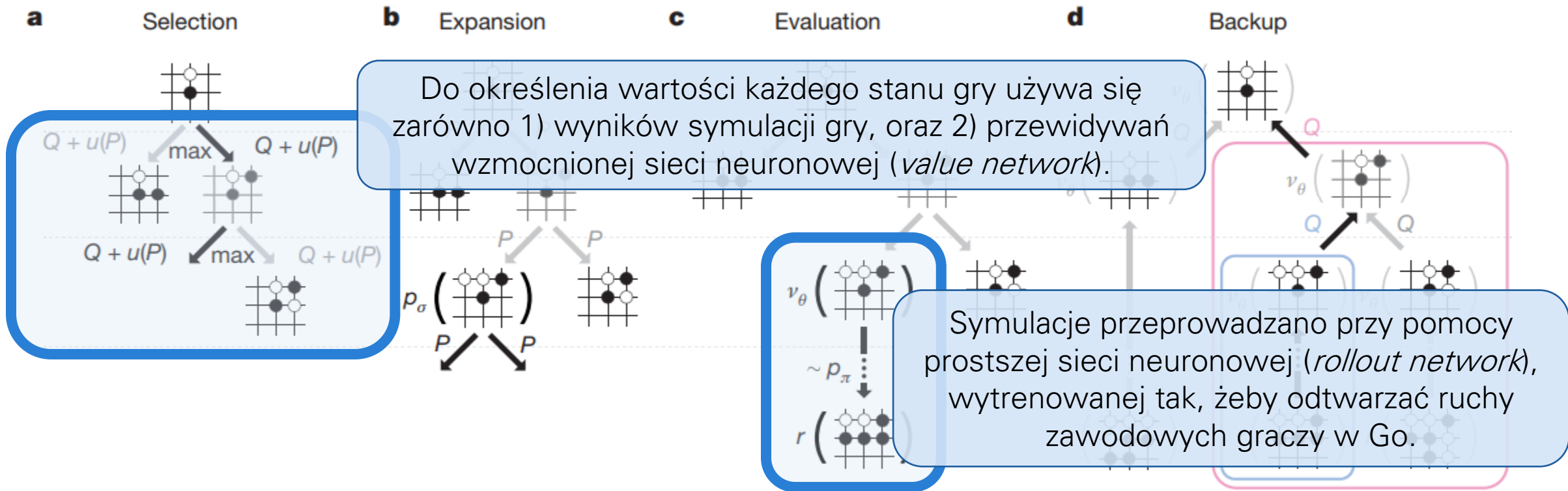
Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: trenowanie sieci neuronowych



Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

Algorytm AlphaGo: MCTS



Źródło: D. Silver, A. Huang, C. Maddison, et al. *Mastering the game of Go with deep neural networks and tree search*. Nature 529, 484–489 (2016).

AlphaGo vs mistrzowie Go

- W październiku 2015 AlphaGo pokonał mistrza w Go, Fan Hui, wygrywając wszystkie 5 rund.
- W marcu 2016 AlphaGo stoczył pojedynek z jednym z czołowych arcymistrzów w Go, Lee Sedol, wygrywając 4 z 5 rund.



AlphaGo - The Movie | Full award-winning documentary



Google DeepMind
538K subscribers

Subscribe

301K



Share

Save



Algorytm AlphaZero

- Do **grudnia 2017** powstała nowa wersja algorytmu, *AlphaGo Zero*, który został wytrenowany bez wykorzystania wiedzy nt. strategii stosowanych przez ekspertów.
- Do **grudnia 2018** na jego bazie został opracowany algorytm *AlphaZero*, który był w stanie grać na poziomie mistrzowskim nie tylko w Go, ale też szachy i shogi (japońską odmianę szachów).

ARTICLE

doi:10.1038/nature24270

Mastering the game of Go without human knowledge

David Silver^{1*}, Julian Schrittwieser^{1*}, Karen Simonyan^{1*}, Ioannis Antonoglou¹, Aja Huang¹, Arthur Guez¹, Thomas Hubert¹, Lucas Baker¹, Matthew Lai¹, Adrian Bolton¹, Yutian Chen¹, Timothy Lillicrap¹, Fan Hui¹, Laurent Sifre¹, George van den Driessche¹, Thore Graepel¹ & Demis Hassabis¹

A long-standing goal of artificial intelligence is an algorithm that learns, *tabula rasa*, superhuman proficiency in challenging domains. Recently, AlphaGo became the first program to defeat a world champion in the game of Go. The tree search in AlphaGo evaluated positions and selected moves using deep neural networks. These networks were trained by supervised learning from human expert moves, and by reinforcement learning from self-play. Here we report an algorithm based solely on reinforcement learning, without human data, guidance or domain knowledge, that achieves superhuman performance, winning 100 games out of 100 against the previous champion.

RESEARCH

COMPUTER SCIENCE

A general reinforcement learning algorithm that masters chess, shogi, and Go through self-play

David Silver^{1,2*,†}, Thomas Hubert^{1*}, Julian Schrittwieser^{1*}, Ioannis Antonoglou¹, Matthew Lai¹, Arthur Guez¹, Marc Lanctot¹, Laurent Sifre¹, Dhharshan Kumaran¹, Thore Graepel¹, Timothy Lillicrap¹, Karen Simonyan¹, Demis Hassabis^{1†}

The game of chess is the longest-studied domain in the history of artificial intelligence. The strongest programs are based on a combination of sophisticated search techniques, domain-specific adaptations, and handcrafted evaluation functions that have been refined by human experts over several decades. By contrast, the AlphaGo Zero program recently achieved superhuman performance in the game of Go by reinforcement learning from self-play. In this paper, we generalize this approach into a single AlphaZero algorithm that can achieve superhuman performance in many challenging games. Starting from random play and given no domain knowledge except the game rules, AlphaZero convincingly defeated a world champion program in the games of chess and shogi (Japanese chess), as well as Go.

Dalsze wyzwania dla sztucznej inteligencji?

STARCRAFT (ALPHASTAR)



TEXAS HOLD `EM POKER (DEEPSTACK)

Twitch.tv/DutchBoy

Score
Dutch.1 18
DeepStack 15

@DutchBoy

ALL IN

DeepStack

4 A

4 A

Q 3 4 2 7

40,000

35,450

Dutch.1

Q 3

Q 3

Restart Match

Hotkeys

Latest Subscriber: **arcadeadam1** Thanks to our Donators: **uwer: \$350.00, Sh**

Hand #	36		
Position	bb		
Player	Stack	Blind	Cards
Dutch.1	35,450	300	Q 3
DeepStack	4,350	150	4 5
Round	Actions	Cards	
P	c R35450 c	Q 3 4 2 7	
F		2 3	
T		7 4	
R			

Result
Dutch.1 wins 40,000 via queens and threes with seven kicker, beating pair of fours with ace, queen, and seven.

Blinds 150/300

UNIVERSITY OF ALBERTA

Podsumowanie

- Początkowo algorytmy były trenowane na bazie wiedzy eksperckiej.
- Sieci neuronowe, wykorzystanie m.in. przez *DeepBlue*, pozwalają wytrenować algorytm do wykonywania ruchu dla dowolnego stanu gry.
- Algorytm połączenie sieci neuronowych z algorytmem MCTS w *AlphaGo*, pozwoliło na dostarczenie bardziej trafnych przewidywań wyniku gry od samych sieci neuronowych.
- Obecnie algorytmy (np. *AlphaZero*) nie wymagają wiedzy eksperckiej – trenują rozgrywając rozgrywki z samym sobą wykorzystując tzw. uczenie wzmocnione.

