

Wstęp do modelowania procesów sensorycznych i percepcyjnych

Konrad Miazga

Poznan University of Technology

Poznań, 2017

Prowadzący

- mgr inż. Iwo Błądek (*iwo.bladek@cs.put.poznan.pl*)
- mgr inż. Konrad Miazga (*konrad.miazga@cs.put.poznan.pl*)

Plan semestru*

- ① ACT-R (Konrad Miazga)
- ② Nengo (Iwo Błądek)
- ③ Framsticks (Konrad Miazga)

Plan semestru*

- 1 ACT-R (Konrad Miazga)
- 2 Nengo (Iwo Błądek)
- 3 Framsticks (Konrad Miazga)

* prowadzący zastrzegają sobie prawo do zmian w planie semestru**

Plan semestru*

- 1 ACT-R (Konrad Miazga)
- 2 Nengo (Iwo Błądek)
- 3 Framsticks (Konrad Miazga)

* prowadzący zastrzegają sobie prawo do zmian w planie semestru**

** ale nie jakichś drastycznych

Zasady zaliczenia

Obecności

- Dopuszczalna są dwie nieobecności nieusprawiedliwione w czasie semestru.
- Każda kolejna nieusprawiedliwiona nieobecność oznacza obniżenie oceny końcowej o pół oceny.
- Znaczne spóźnienie traktowane jest jak połowa nieobecności.
- 30% nieobecności (zarówno usprawiedliwionych, jak i nieusprawiedliwionych) stanowi podstawę do niezaliczenia przedmiotu.
- Usprawiedliwienia (lekarskie) należy dostarczyć w ciągu dwóch tygodni od nieobecności (później nieobecność będzie uważana za nieusprawiedliwioną).

Zasady zaliczenia

Oceny

- Ocena końcowa liczona jest na podstawie punktów uzyskanych w trakcie semestru (minus nieobecności).
- Aby zaliczyć przedmiot, należy uzyskać łącznie co najmniej 50% punktów (w tym co najmniej 50% z każdego projektu).
- Planowane są trzy projekty:
 - ACT-R
 - Nengo
 - Krótki referat na temat wybranej publikacji związanej z przedmiotem

Czym jest modelowanie?

- Black Box



Czym jest modelowanie?

- Black Box



- White Box



Czym jest modelowanie?

- Black Box



- White Box



- Grey Box



Po co tworzyć modele?

- Można modelować to, do czego brakuje teorii.
- Modelami można manipulować lepiej niż ludźmi.
- Modele można precyzyjniej analizować.

Po co tworzyć modele?

Można modelować to, do czego brakuje teorii

- Teorie dostarczają wysokopoziomowego opisu zjawisk, modele nie mogą uciekać od szczegółów (muszą być możliwe do zaimplementowania).
- Uczenie maszynowe pozwala stworzyć działające modele procesów, do których nie posiadamy żadnej teorii (rozpoznawanie emocji?).
- Takie modele — choć nie muszą koniecznie być zgodne z rzeczywistym modelowanym procesem — mogą dostarczyć nam intuicji odnośnie tego jaka może (ale nie musi) być natura procesu.
- Argumentami przemawiającymi za poprawnością modelu są np. możliwość przeniesienia modelu na strukturę neuronową, bądź zdolność predykcyjna takiego modelu (pozwala przewidywać wyniki eksperymentów na ludziach).

Po co tworzyć modele?

Modelami można manipulować lepiej niż ludźmi

- ...i nie mówimy tu o technikach manipulacji.
- Możemy badać wpływ zmiany parametrów / uszkodzenia modelu na jego zachowanie.
- Możemy badać zachowanie modelu w nietypowych środowiskach.

Po co tworzyć modele?

Modele można precyzyjniej analizować

- Możemy zajrzeć „do środka” modelu i zbadać zachowanie wszystkich jego elementów w trakcie działania.
- Przykład: możemy na bieżąco badać stan każdego neuronu w sieci neuronowej (wejścia, wyjścia).
- Przykład: możemy tymczasowo wyłączać pewne elementy modelu aby zbadać ich wpływ na zachowanie modelu.
- Przykład: w głębokich sieciach neuronowych, możemy zajrzeć do każdego neuronu z osobna i zbadać rozpoznawany przez niego wzorzec.

Pozorna złożoność

- Czy jeśli niezrozumiały dla nas proces (*blackbox*) przejawia złożone zachowanie, czy oznacza to, iż reguły jego działania muszą być złożone?

Pozorna złożoność

- Czy jeśli niezrozumiały dla nas proces (*blackbox*) przejawia złożone zachowanie, czy oznacza to, iż reguły jego działania muszą być złożone?
- Proste reguły potrafią prowadzić do złożonych zachowań (przykład: automaty komórkowe)

Pozorna złożoność

- Czy jeśli niezrozumiały dla nas proces (*blackbox*) przejawia złożone zachowanie, czy oznacza to, iż reguły jego działania muszą być złożone?
- Proste reguły potrafią prowadzić do złożonych zachowań (przykład: automaty komórkowe)
- Łatwo przeszacować stopień skomplikowania systemu.

Pozorna złożoność

- Czy jeśli niezrozumiały dla nas proces (*blackbox*) przejawia złożone zachowanie, czy oznacza to, iż reguły jego działania muszą być złożone?
- Proste reguły potrafią prowadzić do złożonych zachowań (przykład: automaty komórkowe)
- Łatwo przeszacować stopień skomplikowania systemu.
- Brzytwa Ockhama

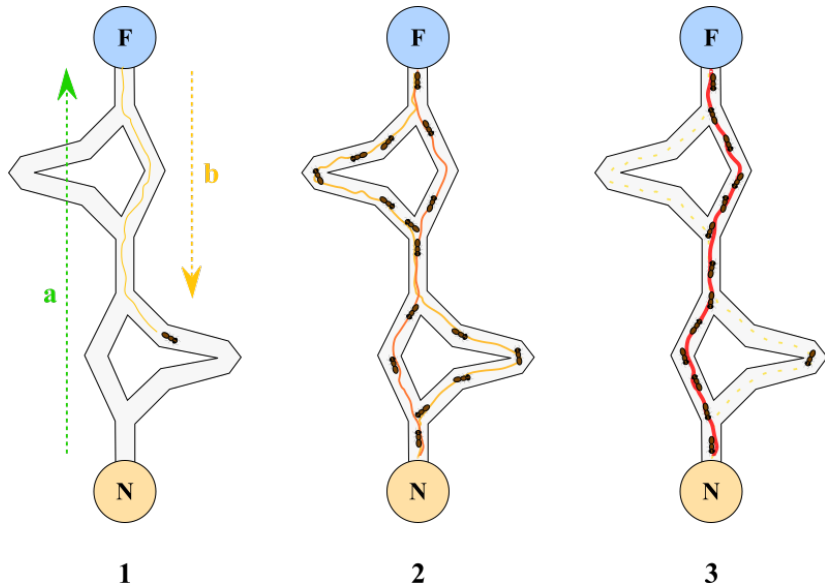
Pozorna złożoność

Przykład #1: Jak mrówki znajdują najkrótszą drogę do pokarmu?

- Mrówki potrafią znajdować krótkie ścieżki do pokarmu.
- W jaki sposób stwierdzają, która ścieżka jest krótsza?
- W jaki sposób przekazują sobie tę wiedzę?
- Czy poszukując jedzenia mrówki zdradzają przejawy inteligencji?

Pozorna złożoność

Przykład #1: Jak mrówki znajdują najkrótszą drogę do pokarmu?



Pozorna złożoność

Przykład #2: „Szwajcarskie roboty”



Pozorna złożoność

Przykład #2: „Szwajcarskie roboty”



Pozorna złożoność

Przykład #2: „Szwajcarskie roboty”



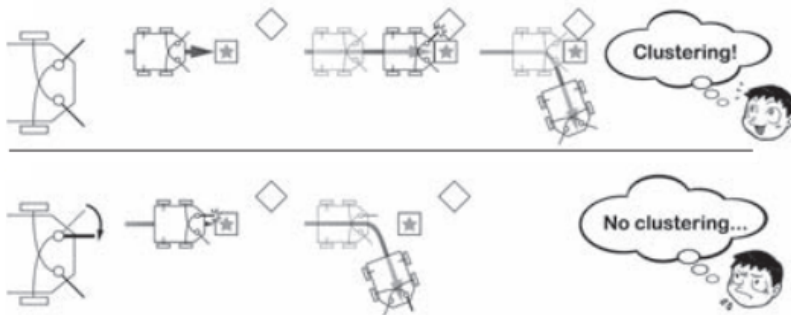
Pozorna złożoność

Przykład #2: „Szwajcarskie roboty”

- <https://www.youtube.com/watch?v=B0wM-eKSxhk>
- Stworzone w latach 90-tych na Uniwersytecie Zuryskim (Maris, te Boekhorst).
- Prosta zasada działania.
- Dwa czujniki pod kątem z przodu robota.
- Robot jedzie do przodu, skręca gdy czujnik wykryje pobliski obiekt.

Pozorna złożoność

Przykład #2: „Szwajcarskie roboty”



Rysunek: *

Źródło: *How the Body Shapes the Way We Think*, Rolf Pfeifer, Josh Bongard

Pozorna złożoność

Przykład #3: Nawigacja mrówki *Cataglyphis*



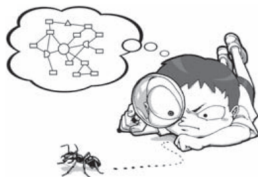
- Saharyjska mrówka pustynna *Cataglyphis* potrafi wracać do domu po samotniczych wycieczkach na odległość nawet do 200 metrów.
- Żyje ona na tunezyjskich solniskach.
- Czy uczy się terenu i tworzy sobie mentalną reprezentację okolicy?

Pozorna złożoność

Przykład #3: Nawigacja mrówki *Cataglyphis*

- Cartright i Collett (1983) sugerują prostszy model.
- Mrówka posiada dwa systemy nawigacji: krótkodystansowy i dalekodystansowy.
- Nawigacja krótkodystansowa: mrówka tworzy *snapshot* horyzontu (ich wzrok obejmuje niemal 360°); gdy wie, że jest blisko domu, kieruje się w stronę maksymalizacji podobieństwa do zapamiętanego widoku.
- Nawigacja długodystansowa: mrówka korzysta z oszacowania odległości do domu oraz kierunku (w oparciu na polaryzację światła słonecznego) aby zgrubnie określić drogę powrotną.

Pozorna złożoność



Pozorna złożoność



Pozorna złożoność



Rysunek: *

Źródło: *How the Body Shapes the Way We Think*, Rolf Pfeifer, Josh Bongard

Podjęcia do modelowania mózgu

- Podejście symboliczne
- Podejście koneksjonistyczne

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne
- W ogólności sekwencyjne

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne
- W pełni rozproszone

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne
- W ogólności sekwencyjne
- Wysoki poziom abstrakcji

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne
- W pełni rozproszone
- Poziom biologicznej implementacji

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne
- W ogólności sekwencyjne
- Wysoki poziom abstrakcji
- Struktury danych i schematy zachowań muszą być ręcznie zdefiniowane

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne
- W pełni rozproszone
- Poziom biologicznej implementacji
- Systemy potrafią się uczyć

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne
- W ogólności sekwencyjne
- Wysoki poziom abstrakcji
- Struktury danych i schematy zachowań muszą być ręcznie zdefiniowane
- Proste do opisania zachowanie i wiedza modelu

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne
- W pełni rozproszone
- Poziom biologicznej implementacji
- Systemy potrafią się uczyć
- Wiedza zawarta w modelu trudna do interpretacji

Podejście symboliczne:

- Związane z tym jak sami postrzegamy nasze działanie
- Oparte na „przesuwaniu symboli”
- Intuicyjne
- W ogólności sekwencyjne
- Wysoki poziom abstrakcji
- Struktury danych i schematy zachowań muszą być ręcznie zdefiniowane
- Proste do opisanie zachowanie i wiedza modelu
- Mało odporne na uszkodzenia systemu

Podejście koneksjonistyczne:

- Związane z tym jak nasze działanie jest widziane okiem neurologa
- Oparte na współdziałaniu wielu agentów implementujących proste obliczenia
- Nieintuicyjne, emergentne
- W pełni rozproszone
- Poziom biologicznej implementacji
- Systemy potrafią się uczyć
- Wiedza zawarta w modelu trudna do interpretacji
- Odporne na uszkodzenia systemu

Podjęcia do modelowania mózgu

Podjęcie symboliczne

- Cele
- Koncepty
- Wiedza
- Percepcja
- Przekonania
- Schematy
- Wnioskowanie
- Akcje

Podjęcie koneksjonistyczne

- Obszary
- Ścieżki
- Neurony
- Membrany
- Jądra
- Kolumny
- Synapsy
- Transmitery

The Central Paradox of Cognition, Smolensky et.al, 1992

"Formal theories of logical reasoning, grammar, and other higher mental faculties compel us to think of the mind as a machine for rule-based manipulation of highly structured arrays of symbols. What we know of the brain compels us to think of human information processing in terms of manipulation of a large unstructured set of numbers, the activity levels of interconnected neurons. Finally, the full richness of human behavior, both in everyday environments and in the controlled environments of the psychological laboratory, seems to defy rule-based description, displaying strong sensitivity to subtle statistical factors in experience, as well as to structural properties of information. To solve the Central Paradox of Cognition is to resolve these contradictions with a unified theory of the organization of the mind, of the brain, of behavior, and of the environment."

Podejścia do modelowania mózgu

Próby symulacji mózgu *in silico*:

- Human Brain Project (Unia Europejska)
- Blue Brain Project (Szwajcaria)
- BRAIN Initiative (Stany Zjednoczone)

Podejścia do modelowania mózgu

Krytyka współczesnych metod badania mózgu:

- “Could a neuroscientist understand a microprocessor?” Eric Jonas, Konrad Kording, 2017

Podjęcia do modelowania mózgu

W kognitywistyce:

Podjęcie symboliczne

- **ACT-R**
- Skupione na reprezentacji procesów mentalnych

Podjęcie koneksjonistyczne

- **Nengo**
- Emergent
- Skupione na implementacji struktury i zachowania mózgu

Podjęcia do modelowania mózgu

W sztucznej inteligencji:

Podjęcie symboliczne

- GOFAI (Good Old Fashioned AI)
- General Problem Solver (Simon, Shaw, Newell, 1959)
- Systemy regułowe
- Skupione na reprezentacji wiedzy

Podjęcie koneksjonistyczne

- ANN (Artificial Neural Networks)
- Deep learning
- Skupione na uczeniu się i adaptacji

Związek między sztuczną inteligencją a modelowaniem mózgu:

- AI – praktyczne, inżynierskie
- Modelowanie – teoretyczne, naukowe
- AI ma na celu efektywne rozwiązywanie problemów
- Modelowanie ma na celu poznanie i realistyczne odwzorowanie rzeczywistości
- AI aktywnie czerpie inspirację z modelowania