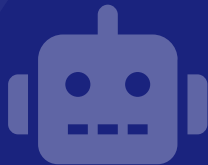


# Artefarita Intelekto

Patxi de la Hoz Méndez (EHU)

84a Hispana Kongreso de Esperanto  
Portugalete, la 16an de Julio, 2026



## ? Kion vi pensas, kiam vi aŭdas “Artefarita Intelektto”?

 Sciencfikciaj robotoj?

 Danĝeraj maŝinoj?

 Io mistera kaj fora?

 Virtuala asistanto?

 Serĉmaŝino?

 Ludo kontraŭ komputilo?

La vero estas **pli simpla** — kaj **pli interesa** ol sciencfikcio!

# Difino: Natura kontraŭ Artefarita Intellekto

## Natura Intellekto (homoj)

- ▶ Lernas el sperto
- ▶ Komprenas kuntekston
- ▶ Havas emociojn
- ▶ Adaptiĝas al novaj situacioj
- ▶ Konscias pri si mem

## Artefarita Intellekto

- ▶ Lernas el **datumoj**
- ▶ Rekonas **ŝablonojn**
- ▶ *Simulas* komprenon
- ▶ Adaptiĝas per **algoritmoj**
- ▶ Ne konscias (aktuale!)

## Difino (simpligita)

**Artefarita Intellekto** (AI) estas komputila sistemo, kiu plenumas taskojn, kiuj normale postulas *homan inteligentecon*.

### Mito 1

“

*AI estas  
konscia.*

#### ✓ Realeco:

AI nur prilaboras  
**ŝablonojn** kaj statis-  
tikajn rilatojn. Ĝi ne  
havas konscion aŭ  
memkonscion.

### Mito 2

“

*AI ĉiam  
pravas.*

#### ✓ Realeco:

AI povas **erari** kaj  
**halucinii** — reprodukti  
malĝustajn informojn  
kun plena konfido.

### Mito 3

“

*AI anstataŭigos  
ĉiujn laborojn.*

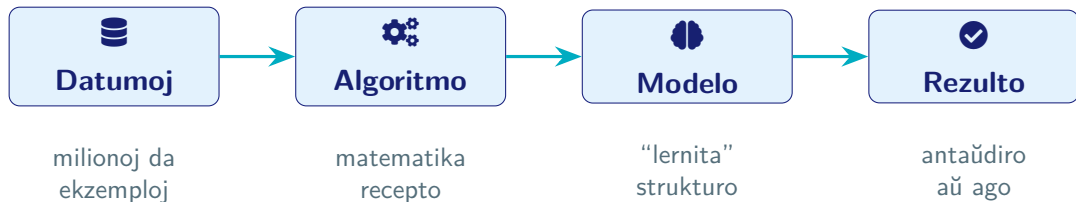
#### ✓ Realeco:

AI plej probable **trans-  
formos** laborojn pli ofte  
ol tute anstataŭigos ilin.

# Mallonga Historio de AI



# La Baza Ideo: Lernado el Datumoj



## Analogio

Infano lernas rekoni katojn vidante **milojn da bildoj** de katoj. Al faras la samon — nur malpli ĝue!

# Tri Manieroj de Lernado



## Gvidata

### (Supervised)

Datumoj *kun* etikedoj.

*Ekz:* “ĉi tiu retmesaĝo estas spamo / ne spamo”

Al lernas la regulon.



## Negvidata

### (Unsupervised)

Datumoj *sen* etikedoj.

*Ekz:* “grupigu similajn klientojn”

Al trovas kaŝitajn ŝablonojn.



## Plifortiga

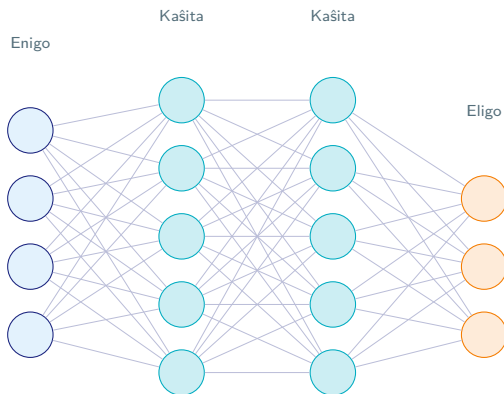
### (Reinforcement)

Al ricevas *premiojn*.

*Ekz:* Al ludas ŝakon — gajno = premio

Al optimumas sian strategion.

# Neŭralaj Retoj — La Cerbo de Moderna AI



## Inspirita de la homa cerbo:

- ▶ **Neŭronoj** = matematikaj funkcioj
- ▶ **Sinapsoj** = pezoj
- ▶ **Lernado** = ĝustigo de pezoj
- ▶ Povas havi *miliardojn* da parametroj

ChatGPT-3 uzis modelon kun  
~**175 miliardoj** da parametroj!

# Mallarĝa vs. Ĝenerala AI

## 🎯 Mallarĝa AI (ANI)

*Artificial Narrow Intelligence*

- ▶ Specialigita por **unu tasko**
- ▶ Povas superi homojn en tiu tasko
- ▶ Ne transiras al aliaj domajnoj

✅ Ĉio ekzistanta AI hodiaŭ

## 🧠 Ĝenerala AI (AGI)

*Artificial General Intelligence*

- ▶ Lernas **iun ajn** taskon
- ▶ Adaptiĝas kiel homo
- ▶ Posedas “komunan saĝon”

⌚ Ankoraŭ ne ekzistas —  
esplorkampo

## ⚠ Grava Distingo

Kiam gazetoj parolas pri AI “pensi” aŭ “konscienci”, ili **troigas**. Hodiaŭa AI estas **tre potenca ŝablona rekonado** — ne magio, ne konscio.

# Kategorioj de Mallarĝa AI



## **Komputila Vidado**

Rekoni vizaĝojn, bildojn, vidaĵojn



## **Natura Lingvo (NLP)**

Traduki, resumigi, respondi en teksto



## **Parolrekono**

Voĉaj asistantoj, transskribo



## **Kreiva AI**

Muziko, bildoj, videoj el teksto



## **Rekomendaj Sistemoj**

Netflix, Spotify, Amazon



## **Predikta Analizo**

Financa aŭ medicina diagnozo



## **Aŭtonomaj Veturiloj**

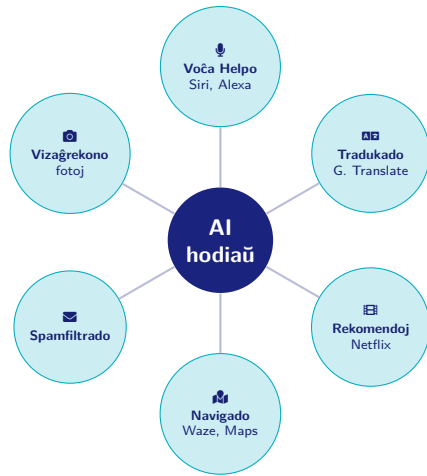
Mem-stirantaj aŭtoj



## **Ludaj AI**

Ŝako, Go, videoludaroj

# AI Estas Jam Ĉie Ĉirkaŭ Vi



## Kio estas LLM?

**Granda Lingva Modelo** (Large Language Model) estas AI trejnita per grandegaj tekstaj aroj (librojn, retpaĝojn, sciencajn artikolojn). Ĝi *antaŭdiras* la sekvan vorton statistike — sed la rezulto aspektas kiel “kompreno”.

### 👍 Utilas por:

- ▶ Skribi kaj kontroli tekstojn
- ▶ Resumi longajn dokumentojn
- ▶ Respondi demandojn
- ▶ Traduki en multajn lingvojn
- ▶ Klarigi teknikajn temojn

### 👎 Limigoj:

- ▶ Povas “halucini” faktojn
- ▶ Sen eksteraj fontoj la modelo ne nepre konas aktualajn informojn
- ▶ Ne komprenas en la sama maniero kiel homoj
- ▶ Povas montri biasojn

# Praktika Ekzemplo: Kiel ChatGPT Funkcias

1. Vi skribas **demandon** (= “prompt”)
2. La sistemo **konvertas** vian tekston al nombroj
3. La neŭrala reto **traktas** la nombrojn tra miliardoj da parametroj
4. La modelo **antaŭdiras** la sekvan tokenon (ofte vorton aŭ parton de vorto) — poste sekvan, sekvan. . .
5. Vi ricevas la **respondon** en teksto

## 💡 Memorinda Metaforo

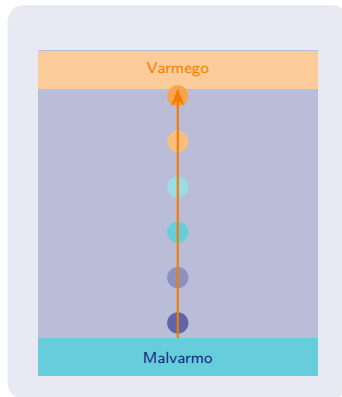
LLM estas kiel **ekstreme klerega aŭtokompleto**. Ĝi scias, kio kutime sekvas, laŭ miliardoj da ekzemploj. Sed ĝi ne “pensas” — ĝi *statistike antaŭdiras*.

## La Sciencia Problemo

Scienco uzas *ekvaciojn* por priskribi la mondon:

- ▶ Kiel varmo disvastiĝas en konstruaĵo
- ▶ Kiel akvo fluas en rivero
- ▶ Kiel aviadilo trairas aeron

Tiuj ekvacioj estas konataj — sed **solvi** ilin por kompleksaj situacioj povas esti *tre malfacile* por tradiciaj komputiloj.

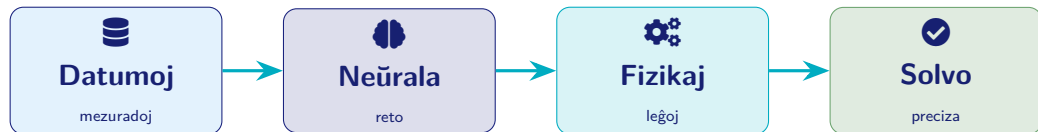


Temperaturaj ŝanĝoj en muro,  
modelitaj de varma ekvacio  $u_t = \Delta u$ .

# Kio estas PINNs?

## Fizike Informitaj Neŭralaj Retoj (PINNs)

**PINN** (Physics-Informed Neural Network) estas neŭrala reto, kiu lernas el datumoj kaj samtempe respektas la *fizikajn leĝojn* de la naturo — ekzemple la leĝojn pri varmo aŭ fluo de fluaĵoj.



# PINNs — Kial Tio Gravas?

## Tradicia metodo:

- ▶ Bezonas **milojn da mezuradoj**
- ▶ Tre **malrapida** komputado
- ▶ Malfacile aplikebla al kompleksaj formoj
- ▶ Postulas specialistajn superkomputilojn

## PINN-metodo:

- ▶ Povas funkcii kun **malmultaj datumoj**
- ▶ Kombinas datumojn kaj fizikajn leĝojn
- ▶ Adaptiĝas al **novaj situacioj**
- ▶ **Promesplena** por kompleksaj problemoj

## ★ Realaj Uzoj

- ▶ **Medicino:** simuli sangofluon en koro
- ▶ **Inĝenierado:** testi pontan strukturon
- ▶ **Klimato:** modeli oceanan varmon
- ▶ **Energio:** optimumigi ventmuelejojn

La graveco: Al ne *anstataŭas* sciencon — ĝi **akcelas** ĝin kaj helpas trakti problemojn antaŭe tre malfacile solveblajn.

# La Bonaj Flankoj de AI

## Medicino

- ▶ Pli frua rekono de kancero
- ▶ Medikamenta malkovro
- ▶ Personigita terapio

## Edukado

- ▶ Persona edukplano por ĉiu
- ▶ Tuja respondo al demandoj
- ▶ Aliro al edukado tutmonde

## Klimato

- ▶ Optimumigo de energiuzo
- ▶ Pli precizaj veterprediktoj
- ▶ Pura energio-distribuido

## Produktado

- ▶ Aŭtomatigo de danĝeraj laboroj
- ▶ Malpli da eraroj en produktado
- ▶ Pli rapida scienca esploro

# La Riskoj — Ni Parolu Honeste

## ⚠ Laboro-Merkataj Riskoj

luj profesioj riskas **aŭtomatigon**:

- ▶ Ĉoferoj, kasistoj
- ▶ Datuma enigo
- ▶ Simpla tekstverkado

*Novaj profesioj ankaŭ aperos!*

## ⚠ Aliaj Riskoj

- ▶ **Deepfakes** — falsaj vidaĵoj
- ▶ **Biaso** en algoritmoj
- ▶ **Privateco** — datuma kolekto
- ▶ **Misinformatado** — AI-generita enhavo
- ▶ **Dependeco** — kritika pensado ↓

La demando ne estas “AI estas bona aŭ malbona?” sed “**Kiel ni uzas AI respondece?**”

# Biaso en AI — Konkreta Ekzemplo

## Kio estas AI-biaso?

Se la **trejnaj datumoj** reflektas antaŭjuĝojn de la socio, la AI **heredas** tiujn antaŭjuĝojn — kaj eble pliigas ilin.

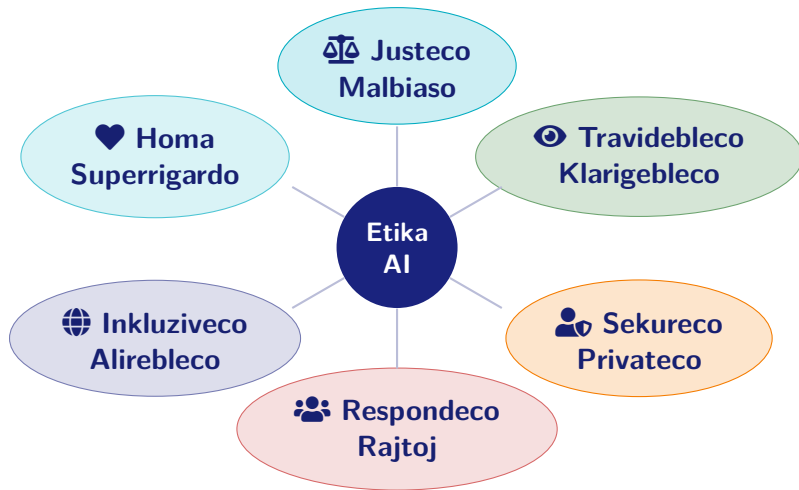
## Realaj ekzemploj:

- ▶ Vizaĝo-rekona sistemo funkciis pli malbone por virinoj kun malhela haŭtkoloro
- ▶ Dungada AI malreklamis virinojn en teknikaj postenoj
- ▶ Kreditaj algoritmoj diskriminaciis laŭ poŝtkodo (= laŭ riĉo)

## 💡 Konkludo

AI ne estas neŭtrala. Ĝi spegulas la datumojn kaj la decidojn de la homoj, kiuj kreis ĝin. **Kritika rigardo estas esenca.**

# Etikaj Principoj por AI



## € Eŭropo

### EU AI-Leĝo (2024)

Klasifikas AI laŭ risknivelo. Alta-riska AI (ekz. medicina) devas plenumi striktajn normojn.

## 🌐 Usono

### Plurfaka Aliro

Federala gvidliniaro, sed ne unu centra leĝo. Industrio reguligas sin parte.

## 🇨🇳 Ĉinio

### Aktiva Reguligo

Striktaj reguloj pri generativa AI. Cenzuro integrita en normojn.

## 🌐 Internacia Debato

Kiel atingi **teknologian progresadon** sen **socia malutilo**? Tio estas la centra debato de nia epoko.

## Baldaŭ (1–5 jaroj):

- ▶ AI-asistantoj en ĉiuj laborlokoj
- ▶ Personigita medicino per AI
- ▶ Pli potencaj tradukaj servoj
- ▶ AI en arto, muziko, filmoj

## Pli malfrue (5–20 jaroj):

- ▶ Aŭtonomaj veturiloj ĝeneraligās
- ▶ Scienca esploro akcelita
- ▶ AI-kuracistaj asistantoj

## ? Malfermaj Demandoj

- ▶ Ĉu AGI estas ebla? Kiam?
- ▶ Kiu posedos la AI-sistemon?
- ▶ Kio okazos al laborlokoj?
- ▶ Kiel ni certigu egalecon?
- ▶ Kiu kontrolas la kontrolistojn?

## 5 Konsiloj por Saĝa AI-Uzado

1. **Kritiku la eligon** — AI-generita teksto povas enhavi erarojn. Ĉiam kontrolu gravajn faktojn en fidindaj fontoj.
2. **Protektu viajn datumojn** — ne enmetu sentemajn personajn informojn en AI-sistemojn, kies privatecaj politikoj vi ne legis.
3. **Lernu kiel uzi AI** — ĝi estas ilo. Bona “prompt” donas pli bonan rezulton. Eksperimentu!
4. **Konservu kritikan pensadon** — ne permesu, ke AI pensu *anstataŭ* vi. Uzu ĝin kiel asistanton, ne kiel oraklon.
5. **Partoprenu la debaton** — AI estas socia fenomeno, ne nur teknika. Via voĉo gravas pri reguligo kaj etiko.

# Ŝlosilvortoj de Hodiaŭ

- ▶ **Artefarita Inteljekto** — komputila sistemo imitanta inteligentecon
- ▶ **Maŝina Lernado** — AI kiu lernas el datumoj
- ▶ **Neŭrala Reto** — AI inspirita de la cerbo
- ▶ **LLM** — Granda Lingva Modelo (ChatGPT, ktp.)
- ▶ **Biaso** — hereditaj antaŭjuĝoj en AI
- ▶ **ANI** — Mallarĝa AI (hodiaŭa)
- ▶ **AGI** — Ĝenerala AI (estonta teorio)
- ▶ **Deepfake** — AI-falsita vidaĵo/voĉo
- ▶ **Prompt** — instrukcioj al AI
- ▶ **Etiko de AI** — respondeca uzado

# La Estonteco de Esperanto en la Epoko de AI

- ▶ Modernaj AI-sistemoj jam kapablas komuniki en Esperanto
- ▶ Ili povas helpi traduki kaj korekti tekstojn
- ▶ Ili povas generi ekzercarojn kaj lernomaterialojn
- ▶ Esperanto havas multe malpli da datumoj ol la angla aŭ la hispana
- ▶ La kvalito daŭre pliboniĝas ĉiujare

## Demando por la estonteco

Ĉu AI helpas malgrandajn lingvojn prosperi, aŭ ĉu ĝi plifortigos la dominadon de grandaj lingvoj?

# Dankon pro la Atento!

Demandoj kaj Diskuto

✉ francisco.delahoz@ehu.eus

| 🌐 <https://sites.google.com/view/dlhm>

